

A PASSIVE HYDROGEN MASER
FREQUENCY STANDARD
DESIGN AND PRELIMINARY RESULTS

THOMAS MATTSSON



LUND 1983

A PASSIVE HYDROGEN MASER
FREQUENCY STANDARD
DESIGN AND PRELIMINARY RESULTS

AKADEMISK AVHANDLING

som för avläggande av teknisk
doktorsexamen vid tekniska
fakulteten vid universitetet i
Lund kommer att offentligen
försvaras å E-sektionen,
hörsal E:B, onsdagen den 25
maj 1983.

kl. 10.15 av

Thomas Mattsson
civ.ing., Mlm

Organization LUND UNIVERSITY Lund Institute of Technology Department of Applied Electronics Box 725 S-220 07 LUND		Document name DOCTORAL DISSERTATION	
		Date of issue 1983 05 25	
		CODEN: LUTEDX/(TETE-1002)/1-148/(1983)	
Author(s) Thomas Mattsson		Sponsoring organization	
Title and subtitle A PASSIVE HYDROGEN MASER FREQUENCY STANDARD DESIGN AND PRELIMINARY RESULTS			
Abstract A <u>passive hydrogen maser frequency standard</u> has been designed by a research group. The frequency standard is based on the transition between the two hyperfine levels ($F=1, m_F=0$) and ($F=0, m_F=0$) in the ground state of atomic hydrogen. A 5 MHz quartz oscillator is frequency locked to the 1420 MHz hydrogen resonance using a phase modulated probe signal. The cavity pulling effect, which is the most important factor limiting the long-term stability of active hydrogen masers, is eliminated in the passively operated maser by a <u>cavity fine tuning servo</u>. The influence of other systematic effects on the accuracy and long term stability of the standard is discussed. The design of the passive maser is not complete in all details, but some preliminary results are presented. The fractional frequency accuracy has been estimated as $1.7 \cdot 10^{-11}$ and a long-term stability of $\sigma_y(\tau) = 2.8 \cdot 10^{-12}$ per day has been obtained through comparison with a cesium frequency standard.			
Key words passive, hydrogen, maser, frequency, standard			
Classification system and/or index terms (if any)			
Supplementary bibliographical information			Language English
ISSN and key title			ISBN
Recipient's notes		Number of pages 164	Price
		Security classification	

Distribution by (name and address) Lund Institute of Technology, Department of Applied Electronics, Box 725, S-220 07 LUND

I, the undersigned, being the copyright owner of the abstract of the above-mentioned dissertation, hereby grant to all reference sources permission to publish and disseminate the abstract of the above-mentioned dissertation.

Signature Thomas Mattsson

Date 1983-04-13

A PASSIVE HYDROGEN MASER
FREQUENCY STANDARD
DESIGN AND PRELIMINARY RESULTS

Thomas Mattsson

Lund Institute of Technology
School of Electrical Engineering
Department of Applied Electronics
Lund Sweden

Lund 1983

LUTEDX/(TETE-1002)/1-148/(1983)



ABSTRACT

A passive hydrogen maser frequency standard has been designed by a research group. The frequency standard is based on the transition between the two hyperfine levels ($F=1, m_F=0$) and ($F=0, m_F=0$) in the ground state of atomic hydrogen. A 5 MHz quartz oscillator is frequency locked to the 1420 MHz hydrogen resonance using a phase modulated probe signal. The cavity pulling effect, which is the most important factor limiting the long-term stability of active hydrogen masers, is eliminated in the passively operated maser by a cavity fine tuning servo. The influence of other systematic effects on the accuracy and long term stability of the standard is discussed.

The design of the passive maser is not complete in all details, but some preliminary results are presented. The fractional frequency accuracy has been estimated as $1.7 \cdot 10^{-11}$ and a long-term stability of $\sigma_y(\tau) = 2.8 \cdot 10^{-12}$ per day has been obtained through comparison with a cesium frequency standard.

T A B L E O F C O N T E N T S

Section	Page
1 INTRODUCTION	1
2 OPERATIONAL PRINCIPLES OF THE HYDROGEN MASER	3
2.1 The hyperfine structure of hydrogen	3
2.2 The active and passive hydrogen masers	9
2.3 Expected performance and limitations	12
2.3.1 Cavity pulling	13
2.3.2 Magnetic field inhomogenities	14
2.3.3 Crampton effect	15
2.3.4 Second order doppler effect	15
2.3.5 Wall shift	15
2.3.6 Spin exchange	16
3 THE DESIGN OF THE PASSIVE HYDROGEN MASER	17
3.1 System overview	17
3.2 Hydrogen Source	22
3.2.1 Palladium leak valve	23
3.2.2 Dissociator	26
3.3 State selector	29
3.3.1 Hexapole magnet	29
3.3.2 Beam geometry	31
3.4 Vacuum system	34
3.5 Storage bulb	35
3.6 Microwave cavity	36
3.6.1 Cavity design	36
3.6.2 Coupling loop design	43
3.6.3 Cavity tuning loop	47
3.6.4 Temperature control	52

Section	Page
3.7 Magnetic field	54
3.7.1 Magnetic field coil	54
3.7.2 Constant current source	59
3.7.3 Magnetic shield	61
3.8 Microwave signal generation	65
3.8.1 5 MHz crystal oscillator	66
3.8.2 5 to 80 MHz multiplier	66
3.8.3 80 to 160 MHz frequency doubler	72
3.8.4 160 to 1440 MHz SRD multiplier	76
3.8.5 Mixing to the hydrogen line frequency	86
3.9 Cavity and hydrogen line servos	89
3.9.1 Modulation frequency generation	90
3.9.2 Phase modulator	93
3.9.3 19.6 MHz AGC amplifier	94
3.9.4 Amplitude detectors	94
3.9.5 Hydrogen line servo	95
3.9.6 Cavity servo	96
4 FREQUENCY DIVIDER IC FOR TICK GENERATION	97
4.1 Introduction	97
4.2 Functional description	98
4.3 Performance evaluation	101
4.4 Tick generator module	105
5 EXPERIMENTAL MEASUREMENTS	107
5.1 Magnetic field measurements	107
5.2 Hydrogen resonance linewidth measurements	109
5.3 Frequency stability measurements	115

Section	Page
6 CONCLUSIONS AND FURTHER RESEARCH	117
6.1 Estimation of total accuracy	117
6.2 Areas for future improvements	118
APPENDIX	121
1 DETAILED DESCRIPTION OF TICK GENERATOR IC DESIGN	121
1.1 Flip-flop design	121
1.2 Simulation and breadboarding of basic circuits.	124
1.3 Layout generation	128
ACKNOWLEDGEMENTS	149
REFERENCES	150

1 INTRODUCTION

The classical application for precision time and frequency standards is generation of time. It started in the 1920's with quartz oscillators used for timekeeping over shorter periods. The rotation of the earth was still the primary frequency and time reference via astronomical observations. As a result of the efforts to increase the long-term stability of crystal oscillators the atomic frequency standards were developed, where the quartz oscillator is locked to a transition between two hyperfine levels in the energy spectrum of a suitable atom. The high performance of the atomic clocks led in 1967 to the redefinition of the second in terms of the cesium resonance. An international atomic time scale (TAI) was generated as the mean value of a small number of laboratory standards and a larger number of commercial standards.

The most important use of time standards is perhaps in the field of navigation. Time measurements may be converted to length measurements via the speed of light. Since a radio signal propagates at the speed of light (app. $3 \cdot 10^8$ sec), a measurement precision of 1 ns is equivalent to a position accuracy of 0.3 m. Today, atomic frequency standards on board orbiting satellites are used in a global positioning system.

Other examples of applications are radio astronomy (very long base interferometry) and scientific experiments for investigation of relativistic effects.

Time laboratories all over the world are operating and developing different types of atomic frequency standards. The most important types are cesium beam tube devices, hydrogen maser devices, rubidium gas cell devices and rubidium maser devices <1>,<2>,<3>. Studies of fundamental accuracy and stability limitations in parallel with technology improvements lead to new ideas of how to further increase the performance of frequency standards. One of the new concepts was the passive hydrogen maser, presented in 1976 by Walls and Howe <10>.

The main disadvantage of the passive hydrogen maser is the complexity of the electronic circuits required. Therefore it was decided that the design of a passive hydrogen maser should be a suitable research project for a time and frequency group at a department of applied electronics. Such a design includes a great variety of electronic circuits from very low frequency servo amplifiers to high frequency microwave components, and the department already had some experience with earlier designs of sodium and rubidium frequency standards.

The group consists of three Ph.D. students, Bo Olsson, Göran Jönsson and myself, and the project leader Dr. Skotte Mårtensson. The work has been divided between the members of the group in the following way:

- Bo Olsson has been responsible for the electronics of the hydrogen line and cavity tuning control systems.
- Göran Jönsson has been responsible for the hydrogen source and the temperature control of the cavity.
- I have been responsible for the microwave frequency generation, the design of the cavity, the magnetic field generation and the operation of the frequency standard.

The aim of this report is to give a description of the theory and the complete design of the passive hydrogen maser although I have not been responsible for all parts of the design. Brief descriptions of parts designed by Bo Olsson and Göran Jönsson are included to achieve better understanding of the complete system. Detailed reports on these parts will be published later.

2 OPERATIONAL PRINCIPLES OF THE HYDROGEN MASER

The basic concept of all atomic frequency standards is to improve the long term stability of a quartz crystal oscillator by means of frequency locking to a stable external reference without aging effects. Transitions between hyperfine levels in the energy spectrum of certain atoms have proved to be such stable references. The considered atomic resonances correspond to transition frequencies in the microwave region (1 - 10 GHz) and have linewidths in the range from 1 to 100 Hz. The frequency of the hydrogen transition used is 1.42 GHz and the linewidth is approximately 1 Hz. The nominal resonance frequency is shifted, however, due to various perturbations. The total accuracy and stability of the atomic frequency standard are depending on these systematic limitations.

2.1 The hyperfine structure of hydrogen

The energy levels of hydrogen have been known for a long time as the result of spectroscopic measurements. Due to the simple structure of the hydrogen atom (a positive proton and a negative electron) the experimental results have been theoretically verified, originally by Bohr's theory of the atom and later by modern quantum mechanics theory.

Solving the Schrödinger equation for the potential energy of the electron in the spherically symmetric case results in a number of discrete energy levels <4>

$$E_n = - \frac{\hbar^2}{2ma^2} \cdot \frac{1}{n^2} \quad (1)$$

where n = the principal quantum number

$\hbar = h/2\pi$

h = Planck's constant

m = the mass of the electron

$a = \frac{\hbar^2}{m \cdot e^2} \cdot 4\pi \cdot \epsilon_0 =$ the Bohr radius

e = the charge of the electron

ϵ_0 = the dielectric constant in vacuum

The principal quantum number n is a positive integer which leads to a minimum energy of -13.6 eV for $n=1$ and an energy equal to zero when n approaches infinity. When an electron

transits from a state with a high energy to a state with a lower energy, the energy difference is emitted as electromagnetic radiation.

In the spherical symmetric case the electron has no orbital angular momentum. If the principal quantum number is greater than 1, however, other types of rotationally symmetric electron orbits exist, having an orbital angular momentum greater than zero. The absolute value L of the orbital angular momentum $\vec{L}$ can be described as a function of l , the azimuthal quantum number, which is an integer smaller than n .

$$L = \sqrt{l(l+1)} \cdot \hbar \quad (2)$$

A charged particle with a non-zero orbital angular momentum is equivalent to a circular electric current. Thus it represents a magnetic dipole with a dipole momentum $\vec{M}$ depending on azimuthal quantum number l .

$$\vec{M} = - \frac{e}{2m} \cdot \vec{L} \quad (3)$$

$$M = - \sqrt{l(l+1)} \cdot \mu_B \quad (4)$$

$$\text{where } |\mu_B| = \frac{e \cdot \hbar}{2m} = \text{the Bohr magneton}$$

In an external magnetic field the magnetic dipole is oriented in a direction such that the projection of the magnetic dipole momentum M on the direction of the magnetic field is an integral number of Bohr-magnetons. This number is called the magnetic quantum number, m , and can range from $-l$ to $+l$.

An intrinsic angular momentum with a corresponding magnetic momentum is caused by the rotation of the electron. The magnitude of this angular momentum, known as the spin $\vec{S}$, is

$$S = \sqrt{s(s+1)} \cdot \hbar \quad (5)$$

where s is the spin quantum number, which is always equal to $1/2$. The component parallel to an external magnetic field must always be $+1/2 \hbar$ or $-1/2 \hbar$.

The total angular momentum $\vec{J}$ of an electron in an atom is the vector sum of its orbital angular momentum $\vec{L}$ and its spin $\vec{S}$.

$$\vec{J} = \vec{L} + \vec{S} \quad (6)$$

For a given value of the azimuthal quantum number l , there are two values of the quantum number for the total angular momentum j of a single electron.

$$j = 1 + \frac{1}{2} \quad \text{and} \quad j = 1 - \frac{1}{2} \quad (7)$$

The hyperfine structure of the hydrogen atom is determined by the total angular momentum of the atom, $\vec{F}$, which is the vector sum of the total angular momentum of the electron and the angular momentum of the nucleus, $\vec{I}$.

$$\vec{F} = \vec{I} + \vec{J} \quad (8)$$

The magnetic dipole moment of the nucleus, μ_I , is

$$\mu_I = \mu_N \cdot g_I \cdot I \quad (9)$$

where $\mu_N = \frac{m}{M} \cdot \mu_B$ = the nuclear magneton

M = the mass of the proton

g_I = the nuclear g-factor

With a nuclear angular momentum $I=1/2$ for hydrogen and a total angular momentum of the electron $J=1/2$ in the ground state ($n=1, l=0, s=1/2$), the total angular momentum of the atom, F , is either 0 or 1 depending on whether or not the directions of the vectors coincide. In the presence of an external magnetic field the direction of the total magnetic momentum is quantized in the same way as the magnetic moment of the electron. The corresponding magnetic quantum number m will be an integer between $-F$ and F . This results in a splitting of each F -level into $2F+1$ sublevels. The degree of this splitting, called the Zeeman effect, depends on the strength of the external magnetic field B . The energies of the Zeeman sublevels are described by the Breit-Rabi formula <5>

$$W(F, m_F) = - \frac{\Delta W_0}{2(2I + 1)} - g'_I \cdot \mu_B \cdot m_F \cdot B \pm \frac{\Delta W_0}{2} \left(1 + \frac{4m_F}{2I + 1} \cdot x + x^2 \right)^{\frac{1}{2}} \quad (10)$$

$$\text{where } x = -(g_J - g'_I) \cdot \frac{\mu_B}{\Delta W_0} \cdot B$$

$$\Delta W_0 = (I + \frac{1}{2}) \cdot A = h \cdot f_0$$

A = the hyperfine coupling constant

g_J = Lande's g-factor for the electron

g_I = the nuclear g-factor

$$g'_I = \frac{m}{M} \cdot g_I$$

The + sign is used for $F=I+1/2$ and the - sign for $F=I-1/2$.

For the hydrogen case where $I=1/2$ the formula may be rewritten as

$$W(F, m_F) = -\frac{A}{4} - g_I' \mu_B m_F B \pm \frac{A}{2} (1 + 2m_F x + x^2)^{\frac{1}{2}} \quad (11)$$

Four different combinations of F and m_F exist: $(F, m_F) = (0, 0), (1, 1), (1, 0), (1, -1)$. Fig. 1 shows the energies of these four sublevels as a function of the magnetic field.

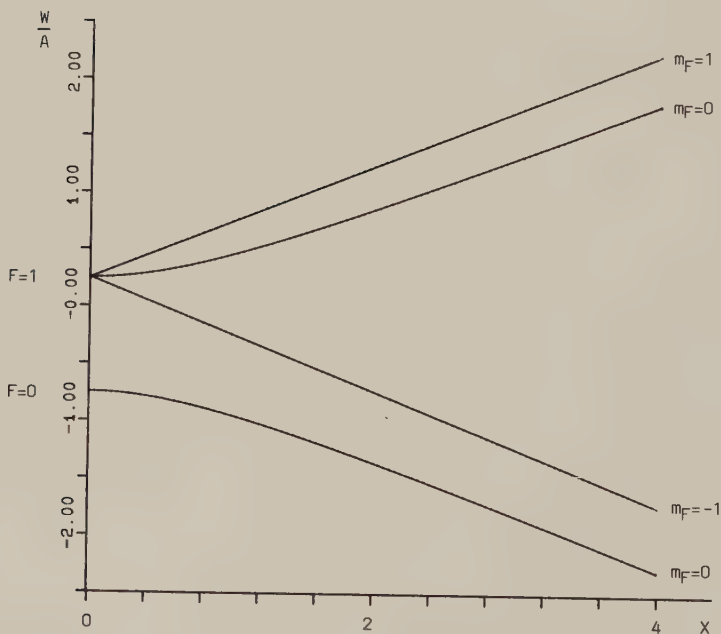


Fig. 1. The hyperfine structure of hydrogen.

Transitions are possible between sublevels with $\Delta F = -1, 0, 1$ and $\Delta m_F = -1, 0, 1$. For low magnetic fields the two levels with $m_F = 0$ show a quadratic dependence on the magnetic field while the other two show a linear dependence. Therefore it is advantageous to use the $(F=1, m_F=0) - (F=0, m_F=0)$ transition because it has lower magnetic field dependence. The expressions for these two Zeeman levels are

$$W(1,0) = -\frac{A}{4} + \frac{A}{2}(1+x^2)^{\frac{1}{2}} \quad (F=1, m_F=0) \quad (12)$$

$$W(0,0) = -\frac{A}{4} - \frac{A}{2}(1+x^2)^{\frac{1}{2}} \quad (F=0, m_F=0) \quad (13)$$

The energy difference as a function of the magnetic field is

$$\Delta W = W(1,0) - W(0,0) = A \cdot (1+x^2)^{\frac{1}{2}} \quad (14)$$

At low magnetic fields the expression may be simplified to

$$\Delta W = A \cdot (1 + \frac{1}{2}x^2) \quad (15)$$

The energy difference may be converted to a transition frequency f .

$$f = f_0 + k \cdot B^2 \quad (16)$$

where f_0 = the transition frequency at zero field

$$k = \frac{x^2}{2 \cdot f_0}$$

The values of the constants used are:

$$\begin{aligned} f_0 &= 1420405751.769 \pm 0.002 \text{ Hz} <14> \\ g_J &= -2.00231924 <6> \\ g_I &= 3.02814 <7> \\ \mu_B &= 9.27409665 \cdot 10^{-24} \text{ J/T} \end{aligned}$$

The frequency of the transition between the two hyperfine levels ($F=1, m_F=0$) and ($F=0, m_F=0$) is

$$f = 1420 \ 405 \ 751.769 + 2750 \cdot 10^8 \cdot B^2 \quad (17)$$

where B is the magnetic field in Tesla.

It can be shown using quantum mechanics that the energy levels of the hydrogen atom are not sharp but have a finite natural breadth. Calculation of the transition probability as a function of the perturbation frequency leads to a resonance phenomenon. As with other resonant circuits, it is possible to define a Q value and a half-power linewidth for the hydrogen line. The natural linewidth is inversely proportional to the interrogation time and is often broadened by external effects like collisions with other atoms and with walls etc.

The contour of the resonance curve is of the Lorentz shape <4> described by the following equation.

$$A = \frac{1}{1 + (k \cdot \Delta f)^2} \quad (18)$$

A = the normalized power
 Δf = the offset frequency
 $k = 2/B$
 B = the half-power width

Fig. 2 shows the relative output amplitude (proportional to $\sqrt{A}$) from a Lorentz shaped resonance curve with a half-width of 2 Hz.

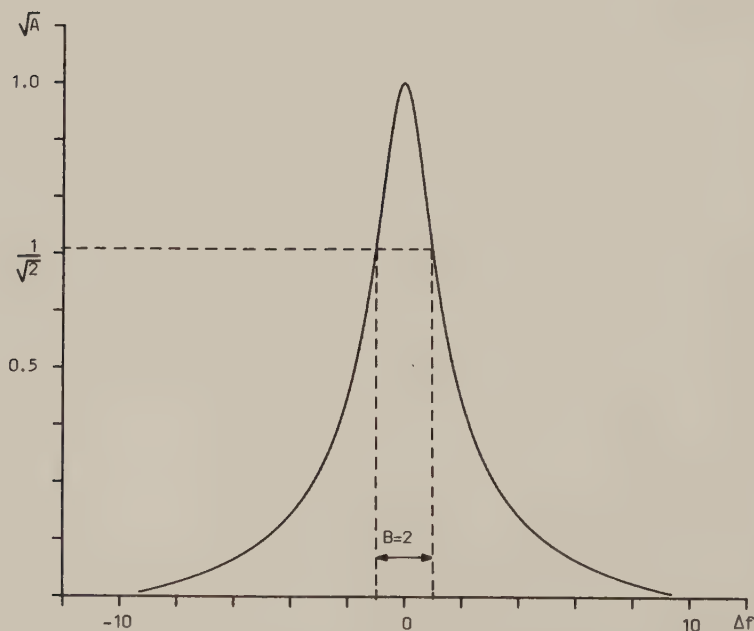


Fig. 2. Output amplitude from Lorentz resonance.

The hyperfine splitting of the $F=1$ level is small at the weak magnetic fields where the frequency standard is intended to operate. The frequency for transitions between these sublevels is in the range from 1 to 10 kHz and can be used to measure the magnetic field as described in chapter 5.

2.2 The active and passive hydrogen masers

The principle of an active maser based on the ($F=1, m_F=0$) - ($F=0, m_F=0$) hyperfine transition at 1420 MHz in the ground state of atomic hydrogen $\langle 8 \rangle, \langle 9 \rangle, \langle 14 \rangle$ is illustrated in fig. 3.

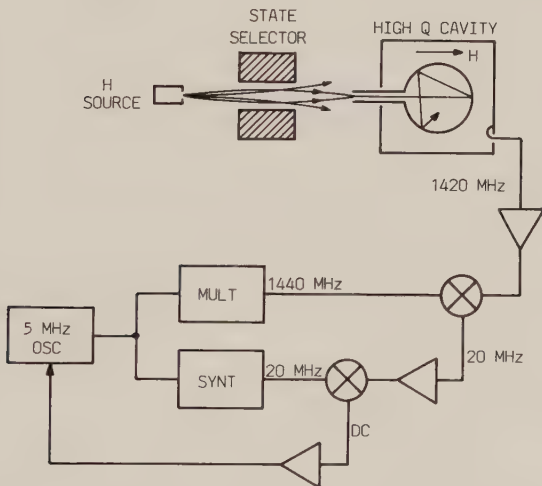


Fig. 3. Block diagram of an active hydrogen maser.

The population difference at thermal equilibrium between the levels of the observed transition is very small due to the small energy difference. Therefore it is necessary to artificially create such a population difference in order to reduce the absorption of energy compared to the emission. This state selection is achieved using the influence of an inhomogeneous magnetic field on the hydrogen atoms. Atoms with a positive effective magnetic dipole moment, i.e. atoms in the ($F=1, m_F=1$) and ($F=1, m_F=0$) states, are focused to the input of the storage bulb by the magnetic field of the state selector. The ($F=1, m_F=-1$) and ($F=0, m_F=0$) atoms with a negative effective magnetic dipole moment are defocused and thus prevented from entering the bulb.

The storage bulb containing the selected atoms is placed inside a high Q microwave cavity tuned to the transition frequency. The electromagnetic field in the cavity stimulates the atoms to transit from the upper level to the lower, emitting their energy at the same frequency and in phase with the existing field. Thus the amplitude inside the cavity increases until the supplied energy from the hydrogen beam is balanced by the losses in the cavity. It is obvious that certain demands regarding the hydrogen beam flux and the Q value of the cavity must be met in order to achieve maser oscillation.

It is important that the resonance frequency of the microwave cavity is the same as the hydrogen transition frequency. Otherwise the output frequency of the maser will be shifted due to cavity pulling. The resonance frequency is kept constant using a temperature compensated cavity design and temperature control of the cavity.

The linewidth of the ensemble of atoms interacting with the microwave field is inversely proportional to the average lifetime of the atoms. To increase the lifetime and thereby decrease the linewidth, relaxation caused by wall collisions is prevented by a teflon coating on the inner surface of the storage bulb. The escape rate from the storage bulb is reduced using a collimator in the inlet tube.

A small uniform magnetic field H inside the cavity is used to separate the hyperfine levels so that only the desired transitions can take place. Due to the quadratic transition frequency dependence on the magnetic field it is, however, desirable to operate at a low field to reduce the frequency instability due to time dependent variations in the magnetic field. The effect of linewidth broadening caused by magnetic field variations within the storage bulb volume is also smaller at lower fields. A lower limit of operation is set by the shielding factor of the magnetic shields that must be used to remove other magnetic fields like the earth's magnetic field.

A small amount of the microwave energy (10^{-12} - 10^{-13} W) is picked up by a small loop inside the cavity and amplified by a low noise amplifier. The output frequency is used to phase lock a 5 MHz crystal oscillator to the hydrogen resonance. Thus the short-term stability of the crystal is combined with the long-term stability of the hydrogen line. The 1420 MHz signal is converted down to 20 MHz in a mixer using a 1440 MHz signal multiplied from the 5 MHz oscillator as a local oscillator frequency. The phase of the 20 MHz signal is compared to the phase of a synthesized 20 MHz signal. The dc error signal from the phase detector is used to adjust the output frequency of the 5 MHz crystal oscillator until it is locked to the hydrogen transition

frequency. The time constant of the servo loop is chosen for an optimum combination of short-term and long-term stability.

The basic principles of the passive hydrogen maser <10>,<11>,<12>,<13> are almost the same as for the active maser. The main difference is the mode of hydrogen atom interrogation. In the active mode a coherent output signal is produced by stimulated emission within a resonant cavity. In the passive mode the hydrogen transition is observed as an amplification of an input microwave signal. The passive mode offers some advantages (which will be discussed later) at the expense of more complex electronic circuitry. The operation of the passive hydrogen maser is illustrated in fig. 4.

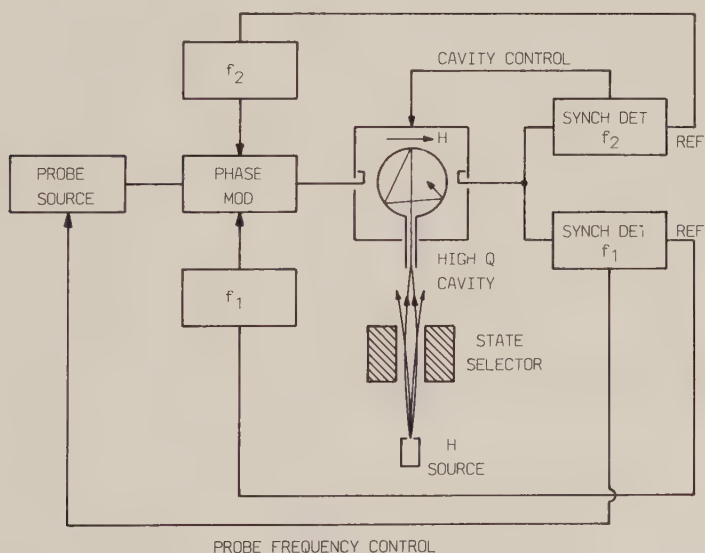


Fig. 4. Principle of the passive hydrogen maser.

The hydrogen source, the state selector, the storage bulb and the uniform magnetic field are exactly the same as for the active maser. The microwave cavity has two additional loops, one for the input signal and one for automatic fine tuning of the cavity. A probe frequency corresponding to the hyperfine transition frequency is phase modulated at two different frequencies, a low frequency f_1 and a high frequency f_2 , before it is fed to the input loop of the

cavity. The low modulation frequency, f_1 , is on the order of the half-width of the hydrogen line and is used to interact with the state selected hydrogen atoms. The higher modulation frequency, f_2 , corresponds to the half-width of the microwave cavity and is used to monitor the resonance frequency of the cavity. The output signal from the cavity is amplitude modulated at the frequencies f_1 and f_2 due to the stimulated emission of the hydrogen atoms and the resonance of the cavity. After envelope detection the f_1 and f_2 signals are compared to their respective reference signals in two synchronous detectors. The phase error of the detected f_1 signal is proportional to the frequency offset between the probe oscillator and the center frequency of the hydrogen line and is used to correct the probe frequency. The phase error of the detected f_2 signal is proportional to the frequency offset between the resonance frequency of the hydrogen line and the resonance frequency of the microwave cavity and is used to correct the fine tuning of the cavity. The time constant of the hydrogen servo must be long to suppress the low f_1 frequency but the time constant of the cavity servo may be shorter.

The greatest advantage of the passive hydrogen maser compared to the active is the ability to control the resonance frequency of the microwave cavity. The cavity pulling effect is probably the most significant factor limiting the long-term stability of hydrogen maser frequency standards. This effect may be almost eliminated using the cavity control technique.

Since no oscillation must be maintained in the passive maser, there are no threshold conditions on the beam flux and the Q of the cavity that have to be met. This offers greater flexibility in optimizing various parameters to achieve better long-term stability.

2.3 Expected performance and limitations

There are a number of systematic effects in the hydrogen maser that affect the output frequency and long-term stability of the passive hydrogen maser.

- 1) Cavity pulling
- 2) Magnetic field inhomogenities
- 3) Crampton effect
- 4) Second order doppler effect
- 5) Wall shift
- 6) Spin exchange

The nature of these effects are discussed in the following sections along with the requirements which must be met to achieve the desired accuracy and long-term stability.

2.3.1 Cavity pulling

Due to the coupling between the hydrogen resonance and the microwave cavity resonance, an offset in the cavity resonance frequency results in a shift of the hydrogen line frequency. The magnitude of the shift is determined by the relation between the cavity Q and hydrogen line Q $\langle 9 \rangle, \langle 10 \rangle$.

$$f = f_0 + \frac{Q_C}{Q_H}(f_C - f_0) \quad (19)$$

where f_0 = the hyperfine transition frequency
 f_C = the resonance frequency of the cavity
 Q_C = the cavity Q
 Q_H = the hydrogen line Q

To obtain a long-term stability better than 10^{-14} with a typical cavity Q of 40000 and a hydrogen line Q of 10^9 , the cavity must be tuned to within 0.3 Hz.

There are a number of time dependent disturbances able to move the resonance of the cavity:

- Cavity temperature
- Bulb temperature
- Mechanical strain
- Electrical changes in the cavity walls
- Changes in the coupling to the cavity

- Leakage of radiation from the cavity to the outside

The influence of these disturbances may be reduced by careful design of the cavity. The cavity may be temperature compensated and temperature controlled to reduce changes in the dimensions. An adequate mechanical design reduces the mechanical strain which may cause deformation of the cavity. Temperature controlled and magnetically shielded isolators should be used at the input and output loops to prevent changes in the reactive load of the cavity. The radiation leaking through openings and slots in the cavity may be absorbed by microwave absorbers so that changes in the environment of the cavity do not affect the resonance.

The flexibility of the passive hydrogen maser allows the design parameters to be chosen in a way that reduces cavity pulling. The hydrogen line Q can be increased by reducing the hydrogen density in the storage bulb and the cavity Q may be decreased. Thus the inherent pulling effect is reduced.

The greatest advantage of the passive maser compared to the active is the possibility to almost eliminate cavity pulling using a cavity servo. If the time constant of the control system is short, the more slowly varying environmental perturbations of the cavity resonance frequency can be compensated by the fine tuning loop.

2.3.2 Magnetic field inhomogenities

The hyperfine transition frequency, f , of the hydrogen maser is related to the magnetic flux density, B , by the relation

$$f = f_0 + 2750 \cdot 10^8 \cdot B^2 \quad (20)$$

In order to keep the frequency offset due to magnetic field deviations within the storage bulb volume less than 10^{-14} , the mean square of these deviations must be kept below $0.7 \cdot 10^{-8}$ T. This is relatively easy to achieve at magnetic fields below 10^{-7} T (it requires a variation of less than 7%). Operating the maser at such a low magnetic field requires a well-designed magnetic shield made of 4-6 layers of μ -metal to reduce the influence of external magnetic fields.

2.3.3 Crampton effect

Another systematic shift of the maser output frequency is caused by the presence of radial rf and dc magnetic fields in the storage bulb region. This effect, called the Crampton effect <17>,<18>, is proportional to the gradient of the radial field and to the population difference in the ($F=1, m_F=1$) and ($F=1, m_F=-1$) levels caused by state selection.

A possible solution to this problem is to increase the magnetic field to $2 \cdot 10^{-6}$ T. This would reduce the effect by a factor of 100, but require extreme stability of the magnetic field which is difficult to achieve. Other ways to reduce the effect are to equalize the populations before the atoms enter the storage bulb or to use double state selection which removes both the (1,1) and (1,-1) states.

Although none of the methods described above is used, the Crampton effect can be reduced sufficiently by using proper magnetic shielding and careful centering of the storage bulb.

2.3.4 Second order doppler effect

The fractional frequency shift Δf induced by the second order Doppler effect is <8>

$$\frac{\Delta f}{f} = - \frac{1}{2} \frac{v^2}{c^2} = - \frac{3}{2} \frac{kT}{m \cdot c^2} \quad (21)$$

where v = the velocity of the atoms
 c = the velocity of light
 k = Boltzmann's constant
 T = the absolute temperature
 m = the mass of the atom

At 313 K the frequency shift is 0.06 Hz with a temperature dependence of $1.3 \cdot 10^{-13}$ per Kelvin. This means that temperature of the storage bulb (and thereby the temperature of the hydrogen atoms) must be kept constant to within 0.07 K in order to achieve a long-term stability of less than 10^{-14} .

2.3.5 Wall shift

The wall shift originates from perturbation of the hydrogen atom on collision with the teflon coated inner surface of the storage bulb. The magnitude of the wall shift is given by <14>

$$f = \frac{W}{D} [1 - a(T - T_0)] \quad (22)$$

where W = the wall shift parameter at T_0
 D = the diameter of the bulb
 a = the temperature coefficient
 T = the absolute temperature

The values of W and a depend on the type of teflon used. The W value for FEP 120 is on the order of (-355 ± 25) mHz cm while a is $(-11.9 \pm 1) \cdot 10^{-3}$ per Kelvin. For a 15 cm bulb the resulting wall shift is a few parts in 10^{11} .

The value of the wall shift parameter is difficult to predict and must be measured using bulbs of different diameters coated in exactly the same way with a given batch of teflon. It is, however, not possible to obtain a better accuracy than 10% $\langle 14 \rangle$, which limits the total accuracy of the hydrogen maser to 10^{-12} . Measurements of the long-term stability of the wall shift indicate a change on the order of parts in 10^{13} per year. The stability may be influenced by a number of factors like the method of coating, bulb operating temperature, surface cleanliness, background pressure stability etc.

2.3.6 Spin exchange

The desire for strong signals results in a fairly high density of hydrogen atoms in the storage bulb. Perturbation of the hyperfine frequency during collisions between atoms causes a broadening of the hydrogen line and a frequency shift $\langle 19 \rangle, \langle 20 \rangle$. Both the line broadening and the frequency shift are proportional to the hydrogen density. The spin exchange contribution to the linewidth is a few tenths of a Hz and the fractional frequency shift is typically parts in 10^{12} for densities used in active masers.

In a passive maser the hydrogen density in the storage bulb may be decreased, thus increasing the line Q and reducing the spin exchange shift by approximately a factor of ten. To achieve a long-term stability of 10^{-14} the hydrogen density must be stabilized to within a few percent. This is accomplished by controlling the pressure in the dissociator bulb and the output power of the rf discharge source.

3 THE DESIGN OF THE PASSIVE HYDROGEN MASER

3.1 System overview

An outline of the passive hydrogen maser system is given by the block diagram in fig. 5. Hydrogen gas from a cylinder of 99.9999 % pure hydrogen is fed through a palladium leak valve into the dissociator. The gas flow is regulated by varying the temperature of the palladium tube so that the pressure inside the dissociator bulb is held constant. The hydrogen gas (H_2) is dissociated into hydrogen atoms (H) by the rf field generated between two capacitor plates on opposite sides of the bulb. The hydrogen atoms travel into the lower high vacuum chamber through a small hole in the dissociator bulb. Because of the low pressure, which gives a long free path, and the shape of the hole, the atoms form a beam. The beam passes through the center hole of the hexapole magnet which acts as a state selector. Atoms of the desired states are focused by the inhomogeneous magnetic field on the collimator in the flange between the upper and lower vacuum chamber. They pass straight through the upper chamber into the storage bulb. Atoms of undesired states are deflected by the state selector and are stopped by the partition wall and pumped away by the lower vacuum pump. The inner surface of the spherical storage bulb is coated by teflon to prevent atoms losing their state when they collide with the walls. Thus the average lifetime of the hydrogen atoms is increased and determined instead by collisions between one another. The storage bulb is placed inside a microwave cavity to allow the atoms to interact with the microwave signal. The small homogeneous magnetic field in the interaction zone is generated by a cylindrical coil outside the cavity fed by a constant current source. The earth's magnetic field is attenuated sufficiently by magnetic shields.

The nominal hydrogen frequency of 1420405752 Hz is derived from a voltage controlled 5 MHz crystal oscillator. The output of the oscillator is buffered in order to be able to drive several 50 ohm loads. The 5 MHz signal that is to be multiplied to the microwave frequency is phase modulated by two different low frequency signals. The 0.5 Hz signal probes the hydrogen line and is used to frequency lock the crystal oscillator. The 12.5 kHz signal probes the resonance frequency of the microwave cavity and tunes it to the hydrogen frequency. The two modulating signals are produced from the 5 MHz signal by a digital divider circuit. The phase modulated signal is then multiplied in three steps up to 1440 MHz. The first stage from 5 MHz to 80 MHz consists of four frequency doublers using double balanced mixers.

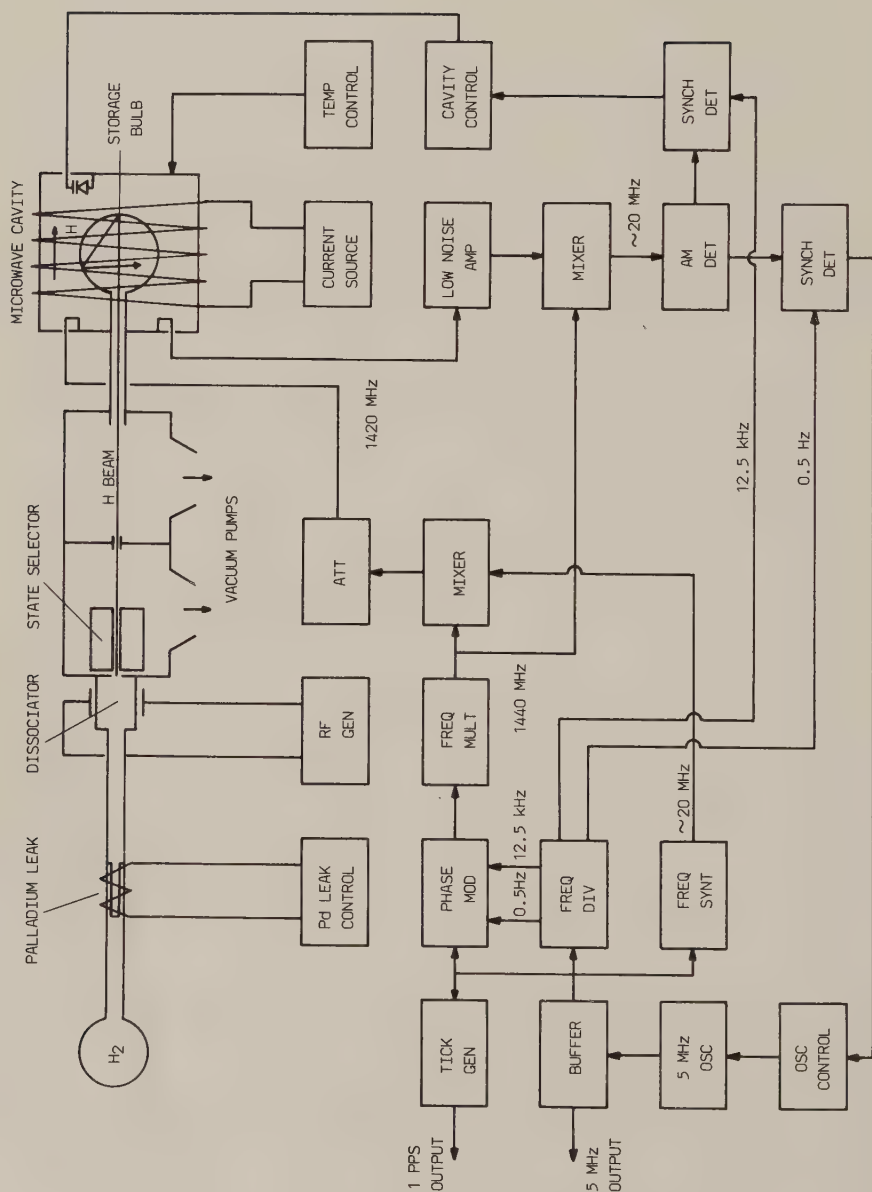


Fig. 5. Block diagram of the passive hydrogen maser

The second stage is a conventional tuned amplifier that doubles the frequency to 160 MHz. The third step is a step recovery diode multiplier that multiplies the signal with a factor of nine to the desired frequency of 1440 MHz. The multiplied signal is then mixed with a synthesized 19594248 Hz signal to obtain a difference frequency that corresponds to the hydrogen frequency. This signal is attenuated to an adequate level and fed into the cavity via the input loop. Inside the cavity the phase modulation generates an amplitude modulation of the microwave signal caused by the stimulated emission of hydrogen atoms and by the cavity itself. Hydrogen atoms which have lost their initial state are removed by the upper vacuum pump.

The output signal from the output loop is amplified in a low noise wideband amplifier and converted to 19594248 Hz in a second mixer using the 1440 MHz signal as local oscillator. After lowpass filtering the 0.5 Hz and 12.5 kHz amplitude modulations are detected in two separate amplitude detectors. The phase of the 12.5 kHz signal is compared with the reference signal in a phase sensitive detector. If the cavity is not correctly tuned there will be a phase difference between the signals. This output signal is used by the cavity servo to produce a control voltage which is put across a varactor tuning diode in a third loop inside the cavity. Thus the cavity is electrically fine tuned to exactly the same frequency as the hydrogen line. Since the fine tune range is limited, the cavity must be held at constant temperature to avoid detuning by room temperature variations. This temperature control is accomplished by varying the current through a heating coil on the outer cylindrical surface of the cavity.

The detected 0.5 Hz signal is compared to its reference in a similar phase sensitive detector. The phase error controls the 5 MHz crystal oscillator frequency by means of a dc voltage from the oscillator servo into the frequency control input of the oscillator. Thus the 5 MHz frequency is adjusted so that the center frequency of the microwave signal coincides with the hydrogen line. The long term stability of the crystal oscillator is now increased and is determined by the hydrogen line instead of the quartz crystal.

The main parts of the passive hydrogen maser are illustrated in fig. 6-8.

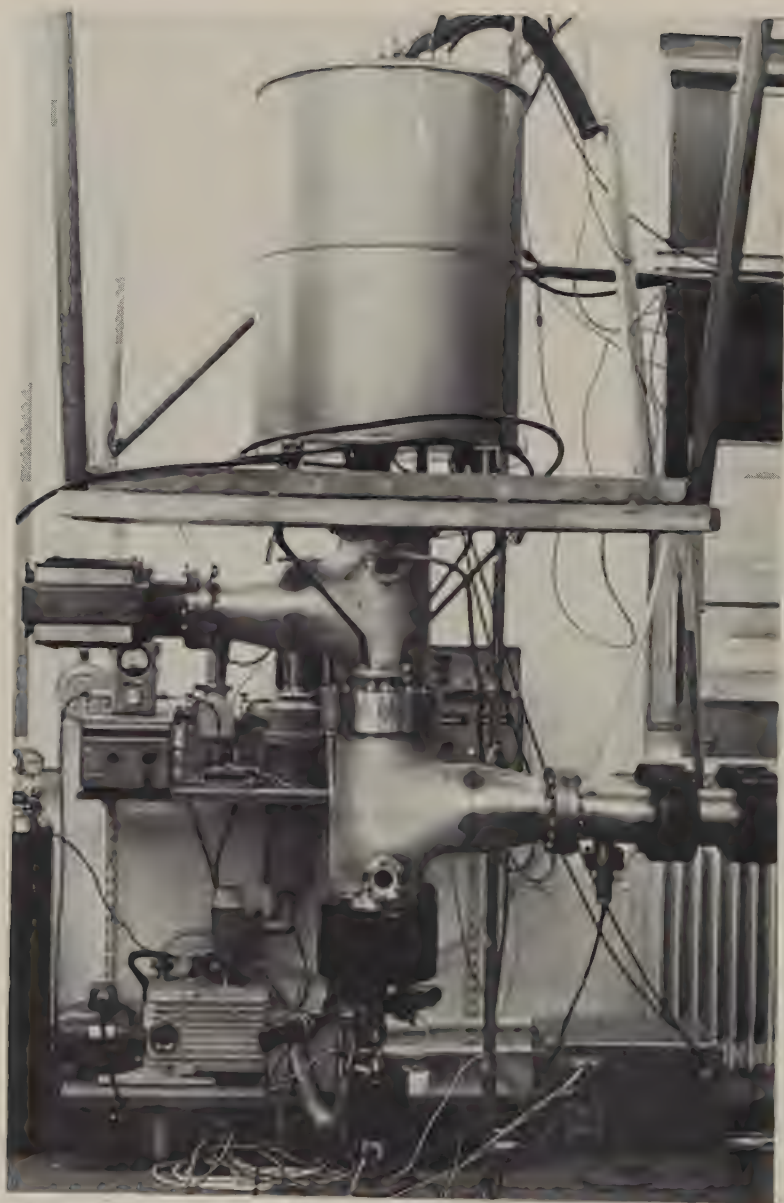


Fig. 6. Photograph of the passive hydrogen maser.

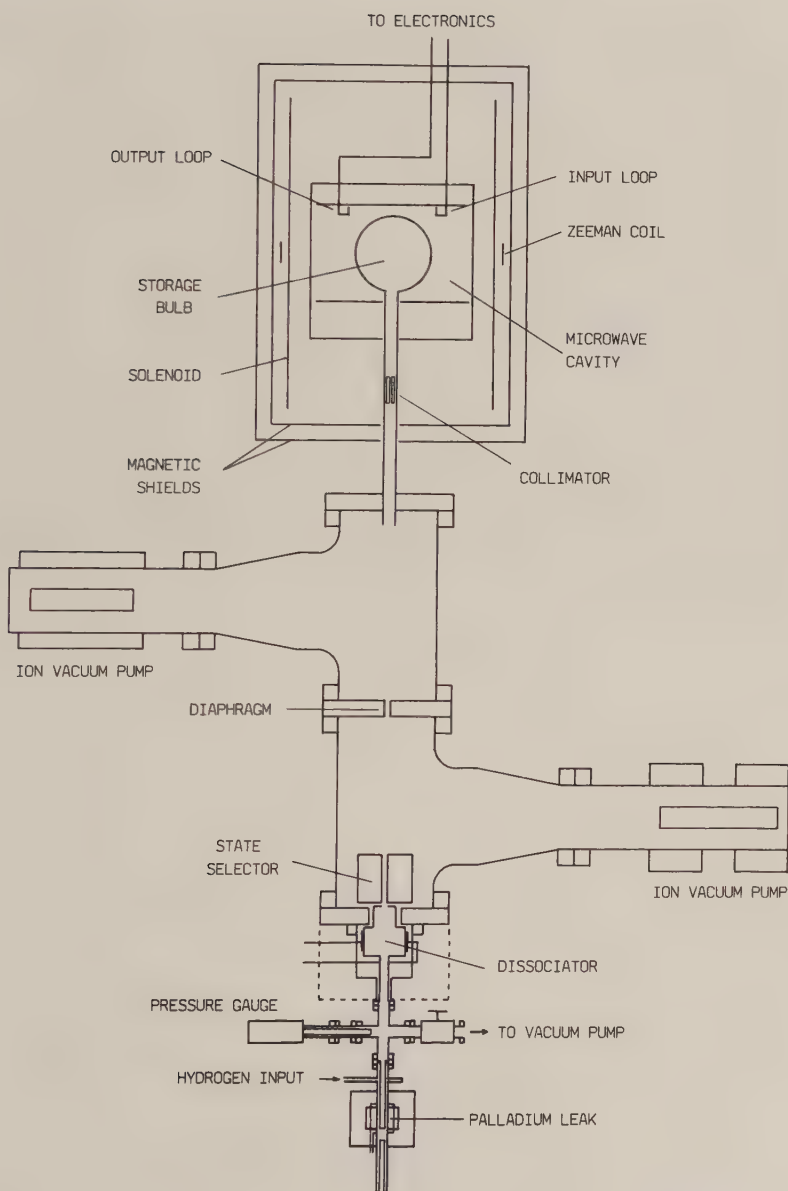


Fig. 7. Schematic diagram of the passive hydrogen maser.

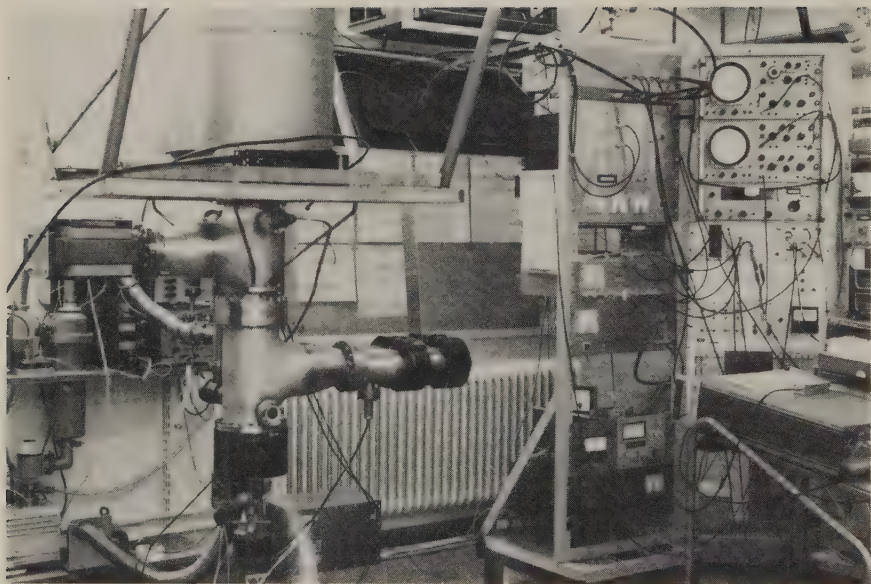


Fig. 8. Photograph of the passive hydrogen maser and the rack mounted electronics.

3.2 Hydrogen Source

The object of the hydrogen source is to produce a beam of hydrogen atoms which is directed into the storage bulb via the state selector. The hydrogen is supplied from some kind of hydrogen gas source, which means that the H_2 molecules must be dissociated into H atoms. It is also necessary to control the intensity of the beam so that it can be changed easily and held constant at different levels.

The dissociation of the hydrogen molecules is done by means of a radio frequency discharge inside a small glass bulb (see fig. 9). The atoms are fed into the high-vacuum chamber via a small hole in the dissociator.

The beam intensity is determined by the size of the hole and the pressure inside the dissociator bulb. This pressure is controlled by an inlet valve that regulates the amount of hydrogen gas flowing into the system.

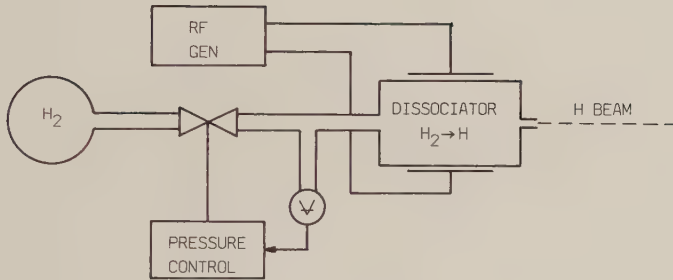


Fig. 9. Principle of the hydrogen source.

3.2.1 Palladium leak valve

Since the flow of hydrogen into the dissociator is comparatively small, it is difficult to use a mechanical type of valve. A better alternative is to use a so-called palladium leak. Palladium has the property that it is permeable to hydrogen. The permeability is dependent on the temperature of the palladium element, which makes it possible to control the flow by varying the operating temperature. To reduce the aging effect at high temperatures an alloy of palladium and silver is used.

In most cases the palladium leak consists of a short length of thin tube that is plugged at one end. Several heating methods can be used to raise the temperature. A method that results in a very short time constant is to let a high dc current flow through the element. Thus the element is heated by the resistive losses in the palladium tube. A simpler method that does not require any wiring inside the low pressure parts is to use an external heater winding. This solution has, however, a much longer time constant which requires a more careful design of the controller.

It was decided that the second method was preferable since the long time constant should not be any problem for the process controller. Fig. 10 shows the mechanical design of the palladium leak.

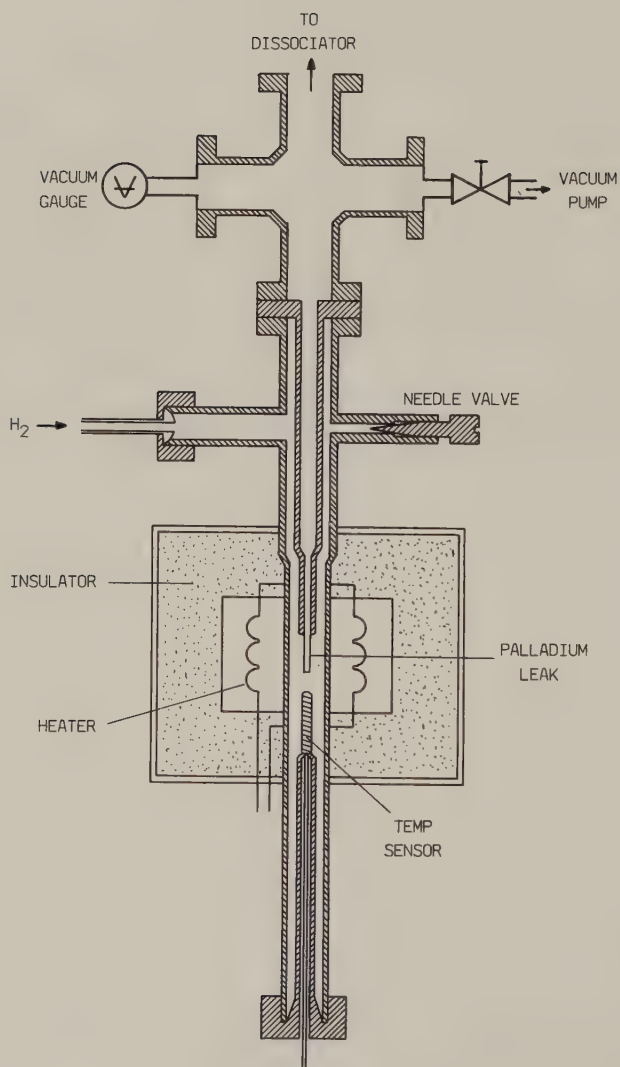


Fig. 10. Mechanical design of palladium leak.

The palladium leak consists of a number of stainless steel

parts mounted together by means of standard high-vacuum flanges. This design is not very compact but it is very flexible and allows one of the parts to be changed without affecting the others.

The hydrogen is supplied from a high pressure gas-holder containing 99.9999% pure hydrogen gas. The pressure is reduced to 0.25 bar by a reducing valve before it is fed to the hydrogen source via a copper tube. The high pressure part of the source can be vented by means of a needle valve to evacuate the air. A small tube, 12 mm in length and 2 mm in diameter, made of 70% palladium and 30% silver is soldered to a concentric inner tube and plugged at the lower end. The palladium tube is heated to the desired temperature by a heater outside the outer tube. The heater is insulated from the surrounding air to reduce heat loss. The temperature inside the heater is measured by means of a Pt100 temperature sensor.

The palladium leak is connected to the dissociator via a tube cross. The pressure in the low pressure part is measured by a standard thermal conductivity vacuum gauge attached to the cross. A valve connected to the rotary vane vacuum pump allows the low pressure part and the dissociator bulb to be evacuated before the source is started.

The pressure inside the dissociator must be held constant to maintain a constant beam intensity. This is accomplished by varying the flow through the palladium leak. During the preliminary experiments the heater has been energized by a constant voltage supply and the pressure has not been controlled. The pressure was monitored on a pressure meter, however, and the variations proved to be small enough to allow meaningful measurements.

In the near future the pressure control in fig. 11 will be implemented.

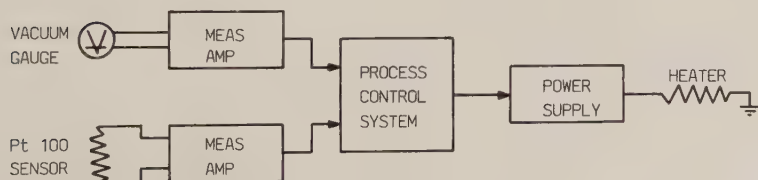


Fig. 11. Pressure control.

The signals from the pressure and temperature sensors are converted by measuring amplifiers to input levels suitable for the process control system. The controller program generates an analog output signal that controls the output power of the heater power supply.

3.2.2 Dissociator

Before the hydrogen leaves the source the molecules must be dissociated so that the beam contains only atoms. The energy required to separate the atoms from each other (4.48 eV) is most conveniently supplied by a radio frequency electromagnetic field. If the energy is higher than the ionization energy (13.6 eV) a gas discharge is started.

If the hydrogen is contained in a small glass bulb, the coupling of the energy to the discharge can be done in either the electric or the magnetic mode. In the electric mode the bulb is placed between the two capacitor plates of a resonance circuit. In the magnetic mode the bulb is placed inside the coil.

Since the dimensions of the first dissociator design were suitable for the magnetic coupling mode, we decided to try this method first. Fig. 12 shows the first dissociator design of a type used in active hydrogen masers at Physikalisch-Technische Bundesanstalt (PTB) in Braunschweig. The cylindrical dissociator bulb with a 0.1 mm diameter collimator was mounted on the bottom flange inside the high vacuum system.

The bulb was surrounded by a three turn coil of 1.5 mm copper wire, tuned to a resonance frequency of approximately 140 MHz by a trim capacitor. The resonance circuit was fed by a tap on the coil. The exact position of the tap was determined by means of admittance measurements. The tap location was adjusted until the input impedance was purely resistive and equal to 50 ohms.

The choice of resonance frequency was influenced by the existence of a suitable power amplifier module designed for that frequency range. A discharge was started at input power levels of 5-10 W. The resonance frequency was almost independent of whether the discharge was started or not and was to a great extent not dependent on the pressure inside the dissociator bulb.

In spite of the good electrical behavior the dissociator did not work properly. The discharge did not assume the expected colour even after several days of continuous operation. However, no leaks or any other explanations of the malfunction could be found. Since we were not too happy

about the fragile mechanical design and the O-ring vacuum seals, we decided to redesign the dissociator.

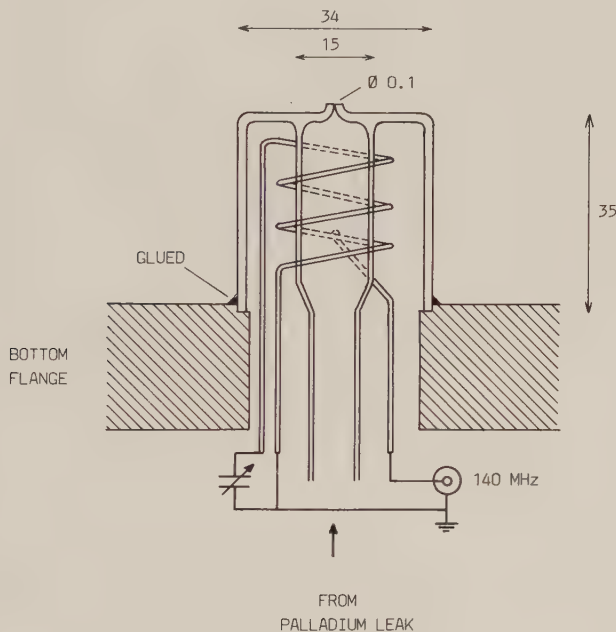


Fig. 12. The PTB-type dissociator design.

A study of the dissociator design of previously built hydrogen masers <12>, <14>, <26> indicated that the volume of the dissociator bulb should be made larger. A new bulb diameter and length of approximately 50 mm was found to be adequate. The diameter of the collimator was at the same time changed to 0.2 mm to increase the beam intensity. The length of the collimator was chosen to 3 mm to achieve a suitable angular distribution of directions of the atoms when leaving the dissociator and entering the state selector.

The large bulb diameter made it difficult to continue to use the magnetic mode coupling without using a lower resonance frequency. Therefore the electric mode coupling was used in the second design.

The new mechanical design (fig. 13) was made more rugged and flexible. All metal parts were made of stainless steel and the O-ring seals were replaced by flanges using copper gaskets.

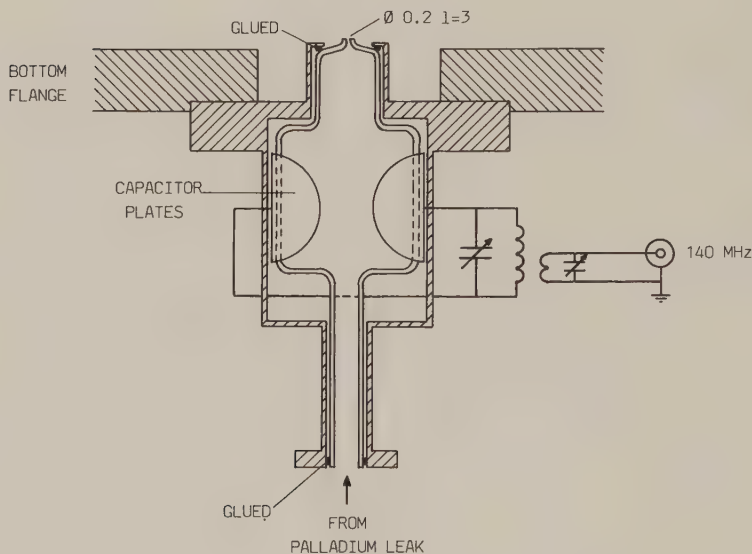


Fig. 13. Improved dissociator design.

Two capacitor plates (30 mm in diameter) for the electric coupling to the discharge are attached to opposite sides of the glass bulb. The plates are connected to the ends of the resonator circuit coil. The resonance frequency is adjusted to approximately 140 MHz by a trim capacitor in parallel with the plates. The power amplifier is transformer coupled to the resonator circuit by means of a second coil. The ratio of transformation is selected for 50 ohm input impedance. The resonance frequency of the electrically coupled dissociator is more sensitive to whether the discharge is started or not than the previously tested design.

The dissociator works for pressures between 0.05 and 0.8 mbar with a minimum input power of 2 W. Normally it is operated at a pressure of 0.4 mbar and 8 W input power. If the resonator circuit is properly tuned the standing wave ratio on the input is better than 1.2. The entire

dissociator is shielded by a perforated iron sheet to prevent rf leakage. The glass bulb is cooled by forced air to reduce the operating temperature.

The efficiency of the dissociation can be evaluated through studies of the light emitted by the discharge. The emitted spectrum is dominated by the α, β and γ lines of the Balmer series. The intensity of the more or less continuous background spectrum, originating from hydrogen molecules and other chemical elements, decreases during the first day of operation to a level far below the discrete lines. This final background intensity should be about 1000 times lower than the intensity of the α line if the dissociation is efficient. Since we only had access to a spectroscope and no possibility to measure the intensities, we could not verify that this criterion was fulfilled.

At present the 140 MHz input signal to the power amplifier module is supplied from a laboratory VHF signal generator. It will later be replaced by an oscillator specially designed for this application.

3.3 State selector

The hydrogen beam that leaves the dissociator contains atoms that can be divided in two groups:

- 1) Atoms with positive effective magnetic dipole moment, i.e. atoms with ($F=1, m_F=1$) and ($F=1, m_F=0$).
- 2) Atoms with negative effective magnetic dipole moment, i.e. atoms with ($F=1, m_F=-1$) and ($F=0, m_F=0$).

The magnetic field from a quadrupole or hexapole magnet has the property that the atoms in group 1 are focused, as opposed to the atoms in group 2 which are defocused. This effect is used to select the desired type of atoms that are sent into the storage bulb.

3.3.1 Hexapole magnet

The hexapole magnet consists of six radially arranged permanent magnets (fig. 14). This magnet configuration generates an inhomogenous magnetic field in the small cylindrical center gap between the poles.

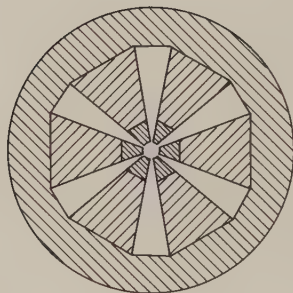


Fig. 14. Cross section of hexapole magnet.

The magnetic field is dependent on the distance r to the center axis.

$$H = H_0 \cdot \left(\frac{2 \cdot r}{D}\right)^2 \quad (23)$$

where D = diameter of the center hole
 H_0 = magnetic field at the pole tips

The magnetic field is zero on the symmetry axis and increases radially towards the poles.

The photograph in fig. 15 shows the hexapole magnet used. It has a 3 mm bore and is 100 mm long.

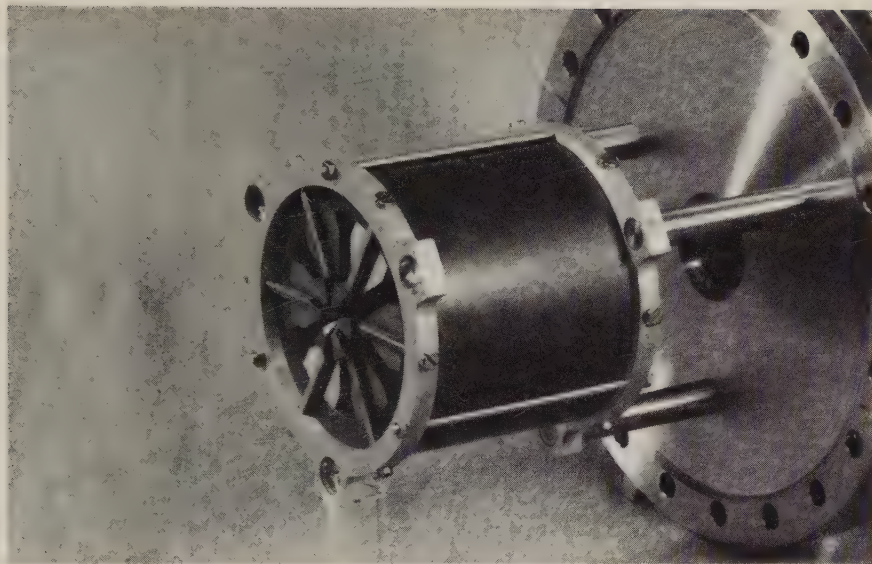


Fig. 15. Hexapole magnet.

3.3.2 Beam geometry

The fundamental beam geometry of the passive hydrogen maser is illustrated in fig. 16.

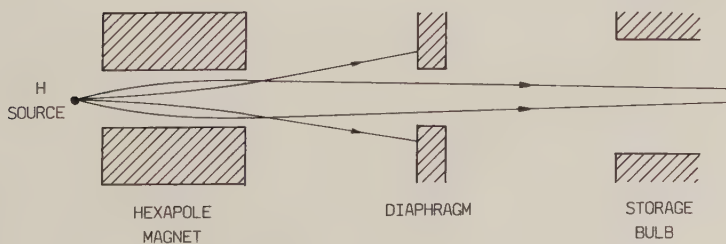


Fig. 16. Fundamental beam geometry.

The hexapole magnet is placed close to the beam source in order to capture most of the hydrogen atoms. Atoms with negative magnetic dipole moment are deflected by the magnetic field and fail to pass through the hole in the diaphragm. In this way these atoms are prevented from entering the storage bulb and are pumped away by the vacuum pump. Atoms with positive magnetic moment, on the other hand, are focused by the state selector magnet.

The path of these atoms through the magnetic field is determined by their entering angle relative to the symmetry axis and their velocity <27>, <28>. If the velocity of the atoms is high the interaction time with the magnetic field is short, and only a small part of the available atoms reaches the target area. The efficiency of the selection, i.e. the ratio between the number of atoms of the right type that enter the hexapole magnet and the number of atoms that pass the diaphragm, increases as the velocity is decreased. Maximum efficiency is achieved for an optimum combination of angular distribution, velocity distribution and aperture dimension.

Since most of the parameters of interest to the beam geometry were unknown at the time the vacuum system was constructed, the dimensions were mainly based on practical considerations. Fig. 17 shows the actual beam geometry.

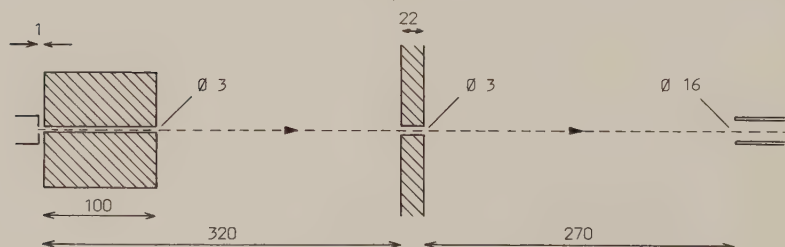


Fig. 17. Actual beam geometry.

The photograph in fig. 18 shows the hydrogen beam source and the hexapole state selector magnet in position on the detached bottom flange of the vacuum system.

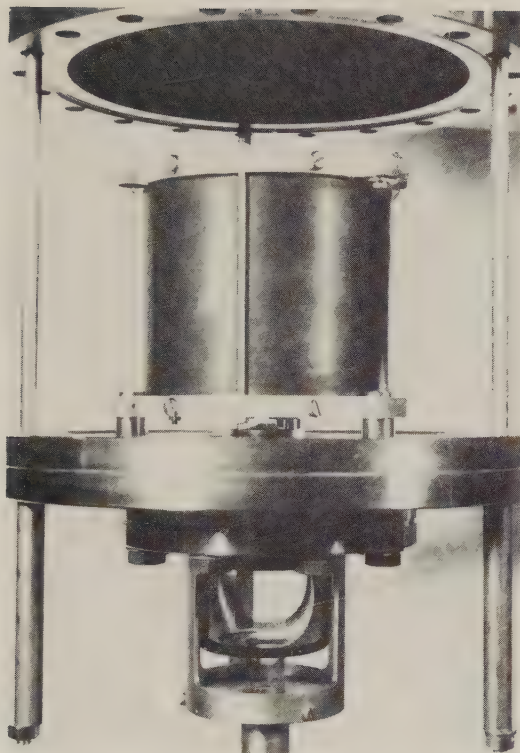


Fig. 18. Beam source and state selector.

Since the magnetic field is zero on the symmetry axis of the hexapole magnet, the atoms passing through the magnet along this axis are not selected. These atoms can be prevented from reaching the storage bulb by means of a beam stop on the center axis. An experiment was made with a beam stop of 0.7 mm diameter located at the exit from the state selector. A reduction of the beam intensity was registered, but no improvement of the hydrogen line width could be noticed. The beam stop was later removed.

3.4 Vacuum system

The pressure inside the hydrogen maser vacuum system must be very low to prevent the beam from being scattered due to collisions with other atoms or molecules. The background pressure in the absence of the hydrogen beam should be lower than 10^{-8} mbar in order to eliminate broadening of the hydrogen line caused by spin exchange collisions with oxygen molecules $\langle 9 \rangle$.

Special requirements concerning materials, construction of components and their treatment have to be met in order to produce such low gas pressures. All parts of the vacuum system are made of stainless steel and joined by flanges with copper seals. All tube dimensions are as large as possible to achieve maximum pump speed.

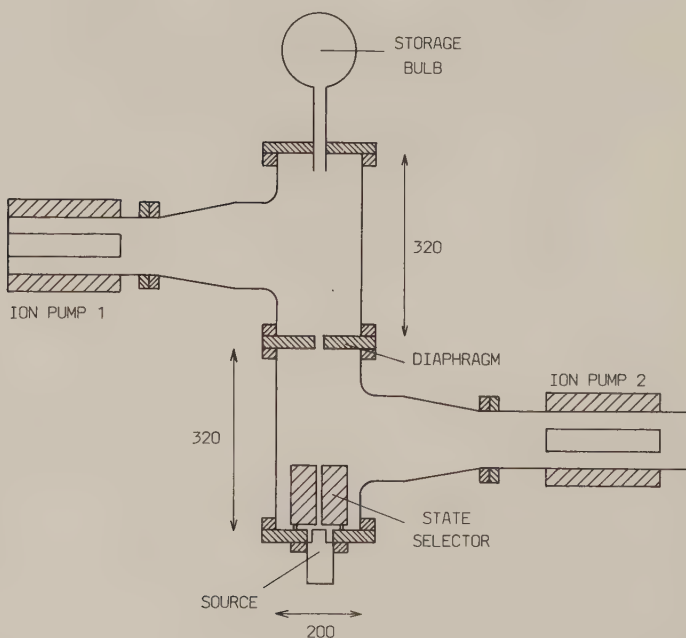


Fig. 19. Vacuum system.

The vacuum system (fig. 19) is divided into two chambers by means of a blank flange with a small aperture. Each chamber has its own ion vacuum pump. Thus the atoms that are

deflected by the state selector and stopped by the diaphragm stay in the lower chamber until they are pumped away by pump 2. In this way the pressure in the upper chamber and the storage bulb can be reduced by a factor of ten compared to the pressure in the lower chamber.

Ion pumps are the most suitable type of vacuum pump for this application. These pumps have high speed, high capacity for hydrogen and do not contaminate the inner surface of the storage bulb. The pumps used have pumping speeds of 80 ltr/sec (upper pump) and 60 ltr/sec (lower pump). Since the ion pumps cannot be started at pressures higher than 10^{-2} mbar, a roughing pump system must be used. The roughing system is connected to the two high-vacuum chambers by means of all-metal valves which are closed when the ion pumps are started.

After several days of baking and outgassing a final pressure of $5 \cdot 10^{-10}$ mbar was reached, which is more than satisfactory. When the hydrogen source is operating the pressure is approximately 10^{-7} mbar in the upper chamber and 10^{-6} in the lower.

It is also recommended that the cavity be evacuated to prevent atmospheric disturbance of the tuning. However, this has not been implemented.

3.5 Storage bulb

The storage bulb used is a spherical bulb made of fused quartz (fig. 20). The diameter is 14 cm and the walls are approximately 2 mm thick. The bulb is connected to the vacuum system by means of a quartz tube with an inner diameter of 16 mm. To allow a more rugged sealing where the tube enters the upper vacuum chamber, the quartz tube is extended by a short length of stainless steel tube of the same diameter.

To prevent relaxation due to collisions with the walls of the storage bulb, the inner surface of the bulb is coated with FEP 120 teflon <29>. The very long teflon molecules on the wall provide more elastic collisions that do not affect the state of the atoms.

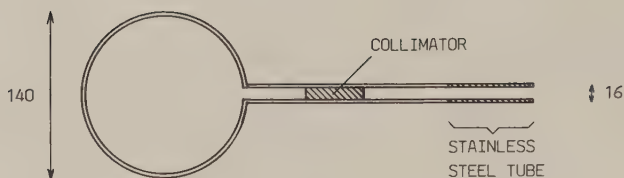


Fig. 20. Storage bulb.

After thorough rinsing of the bulb, the liquid teflon emulsion can be applied until the surface is uniformly wetted. After, drying the bulb is slowly heated to 360°C in an oven to fuse the teflon by melting. Pure oxygen is flowed through the bulb during the procedure in order to oxidize the stabilizer in the teflon emulsion. The gas flow also has the purpose of keeping the teflon surface somewhat cooler than the bulb. Thus the fusing starts from the inside, preventing impurities from being trapped within the teflon. Two additional coatings were applied using the same technique.

To reduce the possibility that the hydrogen atoms will escape the storage bulb, a collimator is placed in the inlet tube in reduce the cross section area. The collimator consists of seven 3.2 mm holes in a 37 mm long cylinder made of solid teflon.

3.6 Microwave cavity

3.6.1 Cavity design

The choice of resonator type has to be a compromise between high Q value and small size. The cylindrical TE_{01n} -mode waveguide resonator offers the smallest volume for a given Q value. Hence the TE_{011} -mode resonance cavity is used in most of the existing hydrogen masers. The dimensions are still quite large however (length = diameter = 30 cm) because of the low operating frequency compared to other standards. Some attempts have been made to reduce the

size of the cavity and thereby the size of the complete maser. Mattison, Vessot and Levine <31> suggest a TE₁₁₁-mode resonator with approximately half the length and diameter. Walls and Howe <12> introduce a dielectric load into the TE₀₁₁ cavity which also reduces the dimensions by a factor of 2. Both types show a lower Q value than the full size cavity but this is not necessarily a disadvantage. Since small size was not of great importance to us for our experimental maser, we decided to use the full size TE₀₁₁ resonator.

For a certain mode the optimum relation between length and diameter can be calculated from the mode-shape factor MS. For TE₀₁₁-mode the mode-shape factor is <30>

$$MS = Q \frac{\delta}{\lambda} = \frac{[x^2 + \pi^2 (\frac{a}{d})^2]^{\frac{3}{2}}}{2\pi [x^2 + (\frac{a}{d})^2 \cdot 2\pi^2]} \quad (24)$$

where δ = skin depth
 λ = wavelength
 a = cavity radius
 d = cavity length
 $x = 3.832$ for TE₀₁

This factor is dependent only on the shape of the cavity (see fig. 21) and has a maximum for $d = 2a$, i.e. when the diameter is equal to the length. The corresponding maximum obtainable theoretical Q value is then

$$Q_0 = 0.659 \cdot \frac{\lambda}{\delta} \quad (25)$$

The skin depth δ depends on which material the resonator is made from. The table below shows calculated Q_0 for a number of possible metals.

Silver	:	81600
Copper	:	79400
Aluminium	:	63400
Brass	:	41200

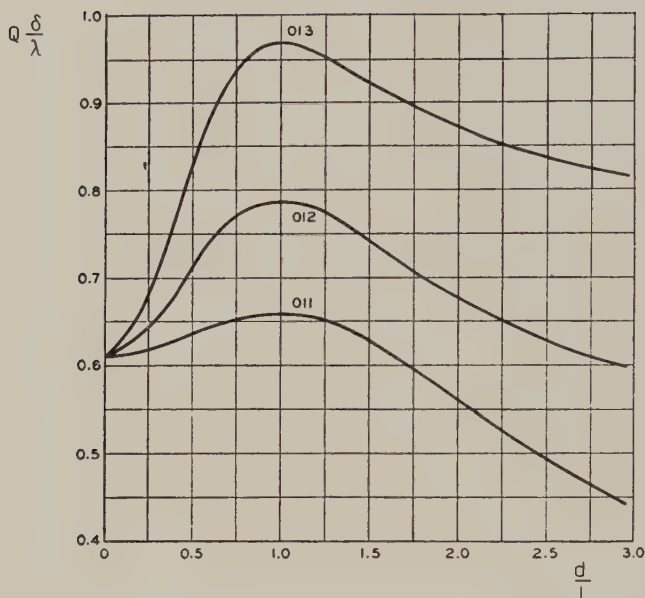


Fig. 21. The mode-shape factor as a function of d/L for the TE_{011} , TE_{012} and TE_{013} modes <30>.

The existence of other resonances close to the desired frequency can be studied in a mode chart (fig. 22). The mode chart is derived from the basic relation between the wavelength in vacuum λ_0 , the cutoff wavelength λ_c and waveguide wavelength λ_g .

$$\frac{1}{\lambda_0^2} = \frac{1}{\lambda_c^2} + \frac{1}{\lambda_g^2} \quad (26)$$

For cylindrical TE_{mnp} or TM_{mnp} waveguide resonators this expression can be rewritten as

$$(f_0 \cdot d)^2 = \left(\frac{c \cdot x}{\pi} \right)^2 + \left(\frac{c \cdot p}{2} \right)^2 \left(\frac{d}{L} \right)^2 \quad (27)$$

where f_0 = the resonance frequency
 d = the diameter of the cavity
 c = the velocity of light
 x = a constant depending on the waveguide mode

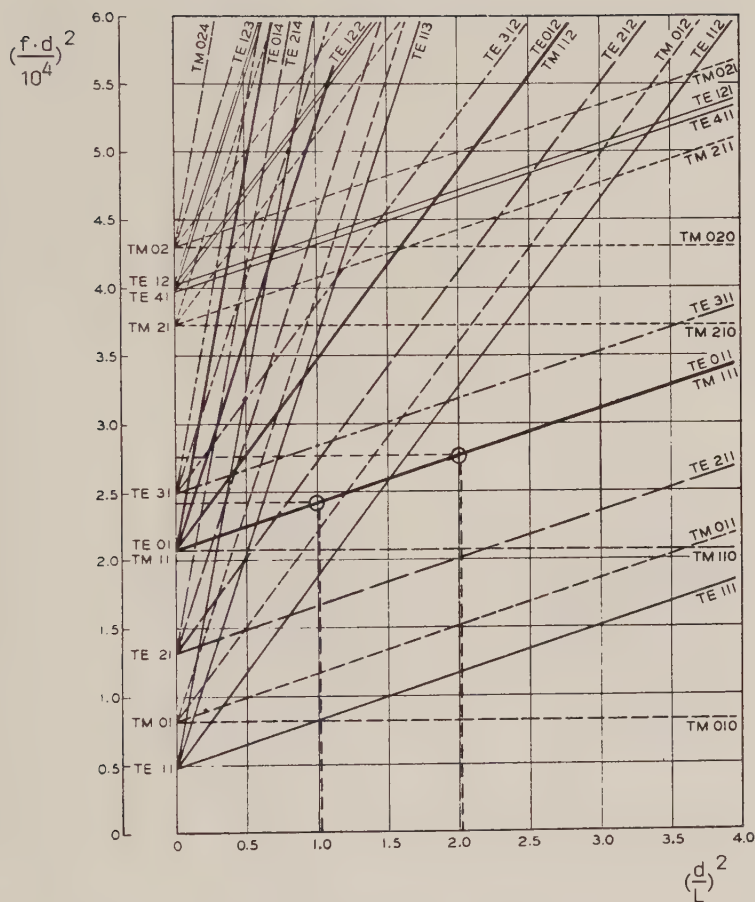


Fig. 22. Mode chart <30>.

Plotting $(d/L) = 1$ for maximum Q value on the diagram shows that there are a number of unwanted resonances in the neighbourhood of the desired operating point. The risk of parasitic resonances is reduced if d/L is increased to a number greater than 1, for example 1.4 ($= \sqrt{2}$). The TM_{111} resonance has the same frequency but a lower Q value and cannot be avoided by changing the cavity dimensions. It has to be eliminated by special measures in the practical design.

A d/L value of 1.4 results in a 2.9 percent lower Q value, which is without significance. The actual dimensions for the 1420405752 Hz resonance frequency were calculated as length $L = 210$ mm and diameter $d = 297$ mm. Minor deviations from these dimensions are allowed during the manufacture and will not have any great affect on performance. Furthermore the length must be mechanically adjustable to enable coarse tuning of the cavity. The calculations above do not consider the influence on the resonance frequency of the storage bulb. The bulb represents a dielectric load inside the cavity and decreases the frequency. To compensate for this, the cavity length must be shortened by approximately 50 mm.

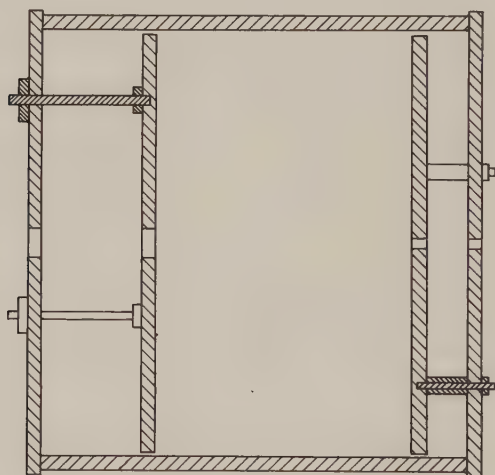


Fig. 23. Cylindrical TE_{011} microwave cavity.

Most hydrogen maser cavities are made of quartz because of its low thermal expansion coefficient. The quartz is silver plated on the inside to obtain a high Q . The temperature dependence of the resonance frequency is often further reduced by mounting one of the end plates on metal rods. Thus the volume change is compensated by a counteracting change in length.

Temperature dependence is less important in a passive maser since a control system can be introduced to fine tune the cavity. Therefore we decided to make it from a material that is easy to tool. A brass alloy was found to be acceptable in spite of its high expansion coefficient.

The cavity consists of a 300 mm long cylinder with two outer and two inner end plates (fig. 23). The outer end plates are just mechanical supports for the inner plates. The inner diameter of the cavity is 297 mm. The upper end plate is mounted with a fixed distance to the outer end plate and contains holes for the input, output and control loops. The lower end plate is mounted on three threaded rods to allow adjustment of the cavity length. It also has a 25 mm hole in the center for the inlet tube to the storage bulb.

The diameter of the inner end plates is made a few tenths of a millimeter smaller than the inner diameter of the cylinder to avoid galvanic contact. Thus the inherent TM_{111} resonance can be suppressed due to the different end plate current distribution for this mode (fig. 24).

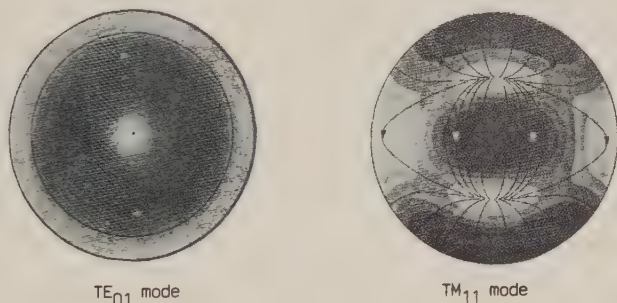


Fig. 24. End plate current distributions $\langle 30 \rangle$.

In the case of TE_{011} there is only a tangential current component (see fig. 25). The current has its maximum at approximately half the radius and is zero in the center and at the periphery.

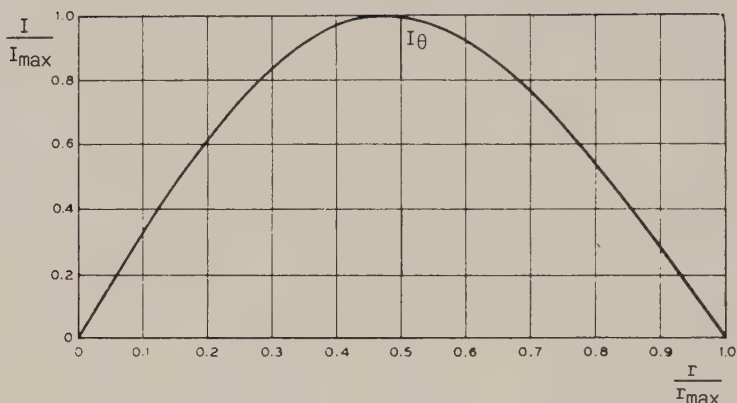


Fig. 25. End plate currents in TE_{01} mode <30>.

The currents in the TM_{111} mode (fig. 26) are a bit more complex. There are both radial and tangential components. The radial component is not zero at the periphery which means that the currents continue in the side wall. If these currents are interrupted by a slit between the end plate and the side wall, an impedance is introduced which suppresses this mode. For the TE_{011} mode the slits are parallel to the current streamlines and there is no such interruption.

It is important that the inner surfaces of the cavity are as smooth as possible. The maximum obtainable Q value is mainly determined by the resistive losses due to the wall currents. The surface resistance not only depends on the resistivity of the metal but also on the distance the current has to flow. If the surface is rough the distance becomes longer and the resistance and losses increase, resulting in a lower Q value.

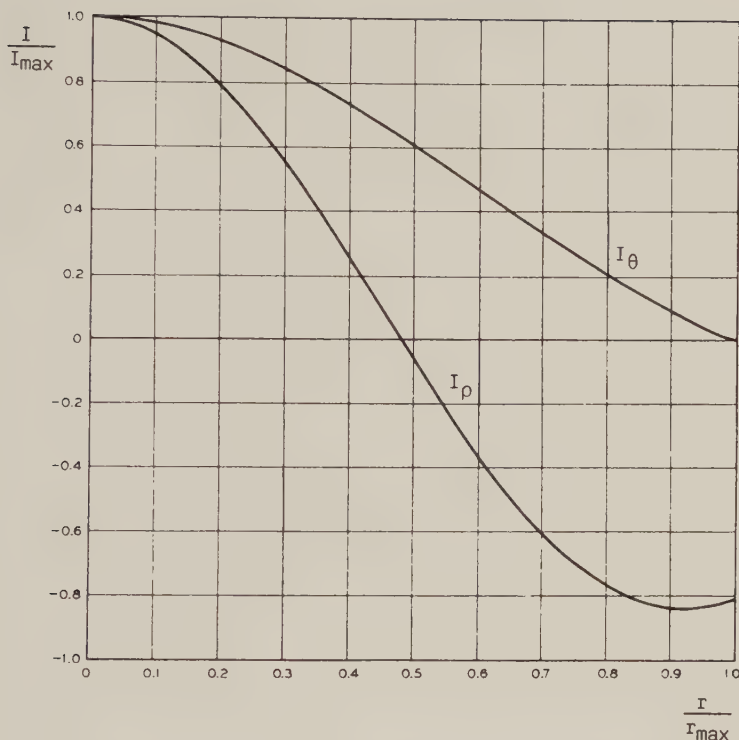


Fig. 26. End plate currents in the TM_{11} mode $<30>$.

3.6.2 Coupling loop design

Coupling from external circuits to the field inside the cavity can be made in many different ways. The two main principles are either to couple to the electric field via a probe or to couple to the magnetic field via a loop. Since the electric field is zero everywhere at the boundary surface of the cavity for the TE_{011} mode, the coupling must be of the magnetic type in this case.

Proper location and orientation of the loop is important to achieve maximum coupling to the desired mode and to contribute to the suppression of all other modes. As shown in fig. 25 the wall currents in the end plates (and thereby the magnetic field) have their maximum at 48% of the radius from the center. Consequently that position would be the best to achieve good coupling. If, in addition to the

magnitude of the field the suppression of unwanted modes is also considered, the optimum distance from the center increases slightly to 54%. Three holes for the input, output and control loops were made at this radius. They were separated as much as possible, i.e. 120 degrees, to avoid direct coupling between the loops.

The first set of loops that were tested had the design shown in fig. 27.

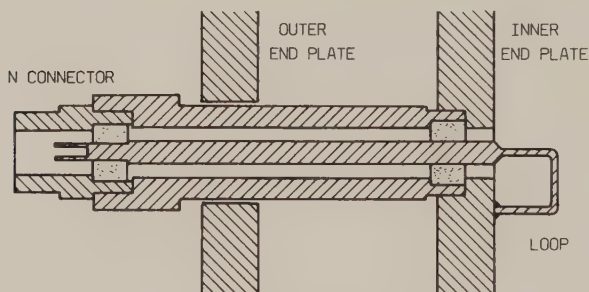


Fig. 27. Original type of coupling loop.

The loop assembly consisted of the following parts:

- A female N-type coaxial connector for the connection of the 50 ohm RG8 coaxial cable.
- A short 50 ohm coaxial transmission line that leads from the connector through the outer end plate to the inner.
- The actual loop inside the cavity.

The dimensions of the transmission line can be calculated from the expression for the characteristic impedance Z_0 as a function of the diameter d of the center conductor, the inner diameter D of the outer conductor and the relative dielectric constant of the insulator in between.

$$Z_0 = \frac{60}{\sqrt{\epsilon}} \ln\left(\frac{D}{d}\right) \quad (28)$$

Since the insulator is air (except for the supports for the inner conductor which are made of teflon) the relative

dielectric constant $\epsilon = 1$. $D = 9$ and $d = 4$ were found to be suitable dimensions for this purpose.

It is important to avoid slits in the cavity walls since they might cause a lower Q value. This is especially true close to the coupling loops due to the high wall current density. Therefore the joints between the end wall and the transmission lines are made inside the end wall. The grounded end of the loops are mounted directly to the end wall. This arrangement has the great disadvantage that the loop cannot be rotated.

The first preliminary measurements of the Q value of the cavity were made with 1 cm^2 loops of the design described above. The results were about ten times lower than the theoretical maximum for a brass cavity without silver coating. Reduction of loop area and wire diameter in the loops gave an improvement of the loaded Q to 12000. This was still far from satisfactory.

The conclusion that could be drawn from these results was that the low Q value did not only depend on the design of the actual loops. The loops are located at places where the wall current density is high. The three 9 mm holes in the end plates where the coaxial transmission lines enter the cavity might cause a severe disturbance of the current streamlines, with low Q value as a result. To reduce this phenomenon a modified type of loop was made.

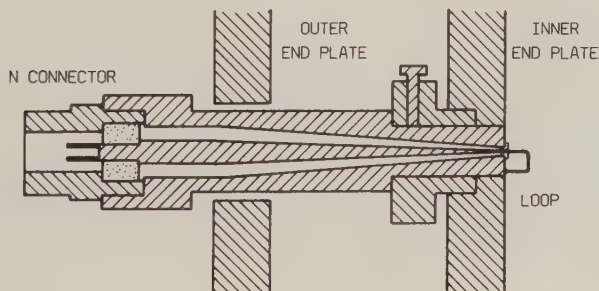


Fig. 28. Modified coupling loop.

The only way to keep the N-type connector and at the same time reduce the diameter of the hole in the end plate, was to use a tapered transmission line. The relation between the conductor diameters is kept constant along the cone to

maintain the 50 ohm characteristic impedance. The measures at the loop end are $D = 2.25$ mm and $d = 1.0$ mm and the new loop area is 0.25 cm^2 (see fig. 29 and 30). Note that the grounded end of the loop is now connected to the outer conductor of the transmission line which makes it possible to rotate the loops, thus allowing the coupling factor to be changed easily.

New measurements confirmed that the assumptions were correct. A loaded Q value of 27000 was achieved and considered to be acceptable. Some additional measurements were made to test the sensitivity to non parallel end plates. Tilting of one of the end plates less than 1 mm had no serious effect on the Q value.

After these adjustments of the cavity design, the inner surfaces were silver plated and protected by a layer of lacquer. The lower resistivity of silver increases the Q value by a factor of 2. The new Q value was measured as 60000.

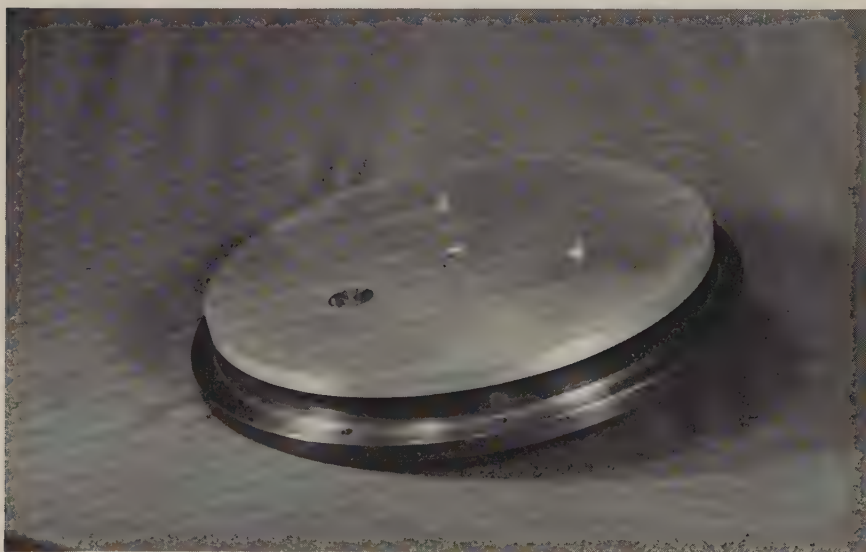


Fig. 29. Photograph of upper end plate with the loops.

So far all measurements had been made with an empty cavity. The influence of the quartz bulb on the length and Q value of the cavity had to be studied. The fact that the relative dielectric constant of quartz is 4.3 causes a distortion of the electromagnetic field inside the cavity. This decreases both the resonance frequency and the Q value. The cavity can

be retuned to the original frequency by shortening its length.

Q measurements after introducing a spherical quartz bulb and reducing the length by approximately 30 mm resulted in values between 45000 and 50000 depending on the position of the bulb.



Fig. 30. Detail of input and output loops.

3.6.3 Cavity tuning loop

One of the main factors that limits the total accuracy of the frequency standard is the cavity pulling effect. Hence fine tuning of the cavity to exactly the same frequency as the hydrogen line is required. In the passively operated maser the resonance frequency of the cavity is probed by a second modulation frequency on the carrier and controlled by a cavity servo system. This method implies a way to fine tune the cavity.

There are basically two different principles of fine tuning:

- 1) Mechanical fine tuning by a metal plunger of variable length inside the cavity.
- 2) Electrical fine tuning by a variable reactive load.

Since there is no obvious choice and both these methods are used in existing masers, it was decided to investigate both

types.

To test the mechanical fine tuning principle, a mechanical fine tuning device was mounted in the center of the upper end plate (fig. 31).

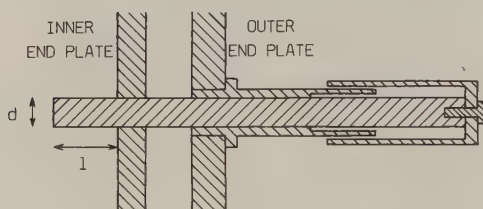


Fig. 31. Mechanical fine tuning device.

A cylindrical metal rod extends into the cavity through a hole in the center of the inner end plate. The length inside the cavity is adjusted by a micrometer screw. The change in center frequency Δf was measured as a function of the length l for two different rod diameters, 4 mm and 10 mm. The diagram in fig. 32 shows the results.

The 10 mm rod gave an acceptable tuning range but the effect of the thinner rod was too small. The frequency detuning is not a linear function of the length ($\Delta f = \text{const} \cdot l^3$). It is however possible to give the plunger a more elaborate shape which gives a more linear behaviour if desired. The Q value of the cavity was only marginally affected.

The first attempt to tune the cavity electrically was made with a reverse biased microwave diode. It was coupled to the cavity via a coaxial wire and a coupling loop identical to the input and output loops. The wire length between the loop and the diode could be varied by a line stretcher. The results of this test were not very encouraging. If the line length was adjusted to maximum tuning effect, the cavity Q was reduced considerably due to losses in the transmission line. If, on the other hand, the length was optimized relative to the Q value, the tuning range was very small. This was obviously not a good solution.

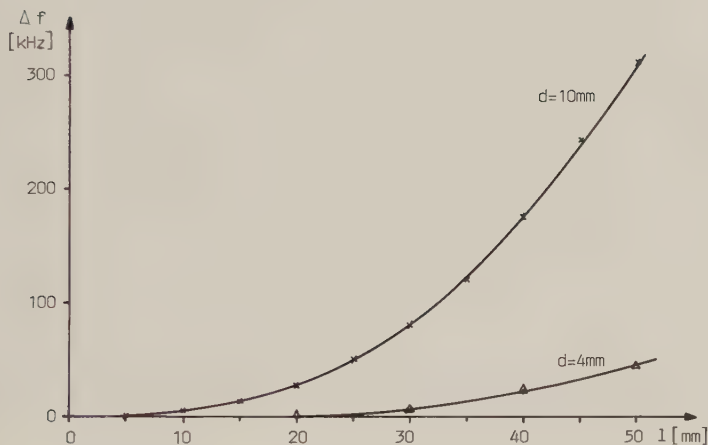


Fig. 32. Frequency change as a function of length for the mechanical tuning.

Another method of tuning a hydrogen maser by means of a variable capacitance diode has been described by Beverini and Vanier <32>. A varactor diode is placed as a part of a loop inside the microwave cavity (fig. 33). This arrangement minimizes the energy losses since no microwave signal is coupled to the outside of the resonator.

A glass encapsulated varactor is soldered at the end of two bars, one grounded and the other the center post of a coaxial arrangement. It consists of two cylindrical capacitor plates separated from ground by a thin dielectric film and a coil on an insulator between the capacitors. This structure a lowpass filter and acts as an rf block on the dc bias supply. The results in fig. 34 were measured using a varactor with the following data:

Capacitance: $C(0\text{ V}) = 5.5\text{ pF}$
 $C(4\text{ V}) = 2.4\text{ pF}$
 $C(45\text{ V}) = 1.0\text{ pF}$
 Q value: $Q(4\text{ V}) > 2000$
 Rev. current: $I_r(36\text{ V}) = 2\text{ nA}$

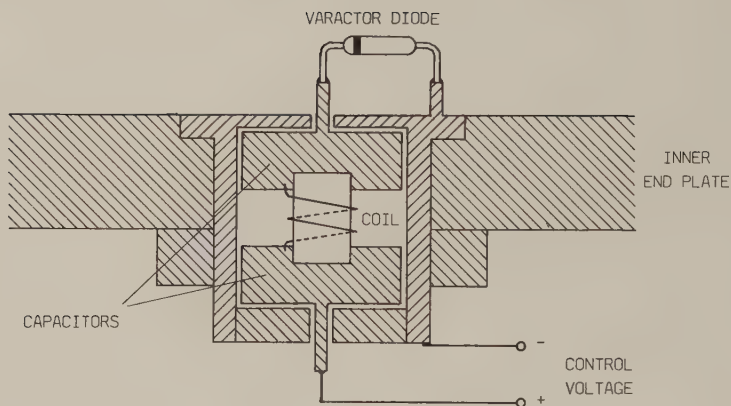


Fig. 33. Varactor tuning loop.

The frequency shows a fairly linear dependence on the varactor voltage and the tuning sensitivity is approximately 2.7 kHz/V. The loaded cavity Q value was only slightly decreased by the tuning loop.

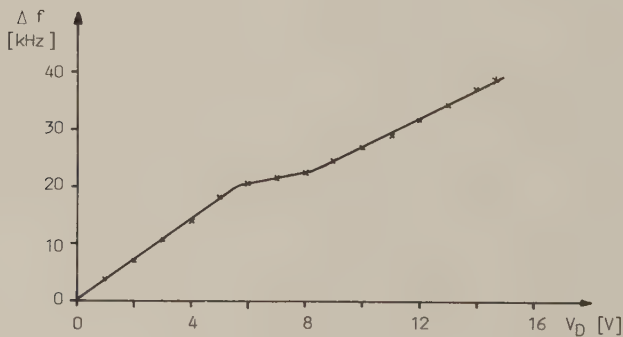


Fig. 34. Frequency offset as a function of varactor voltage.

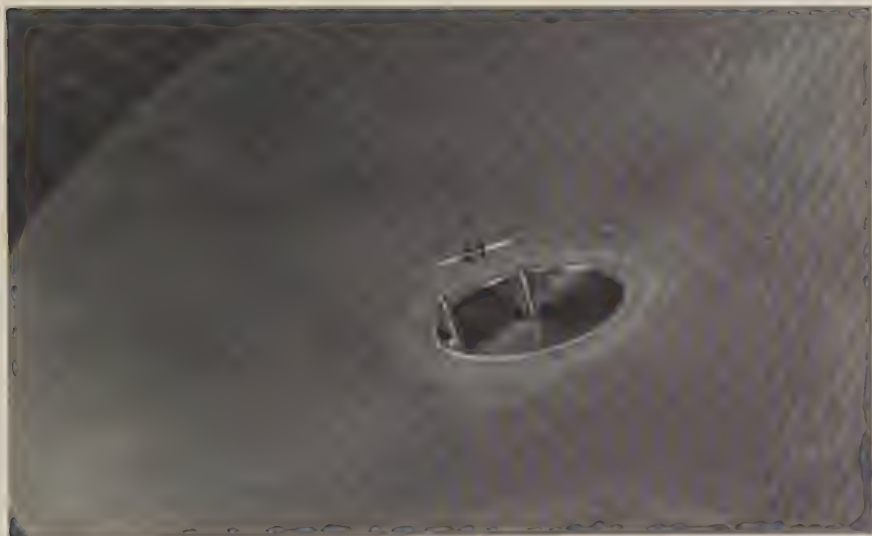


Fig. 35. The mounted cavity tuning loop.



Fig. 36. Detail of cavity tuning loop.

A comparison between the two fine tuning methods shows that they possess almost equal qualities. The mechanical solution has the disadvantage that it requires some kind of stepper motor with associated drive electronics. Therefore the varactor tuning system was chosen and fig. 35 and 36 show the loop mounted in the upper end plate.

3.6.4 Temperature control

As was pointed out earlier, it is important for the total performance of the maser that the microwave cavity is tuned to the hydrogen line frequency. The main reason for changes in the resonance frequency is the change of the physical dimensions of the cavity due to temperature variation. Since the cavity is made of metal, the temperature coefficient of the resonance frequency is large, approximately $27 \text{ kHz}/^{\circ}\text{C}$. To reduce this undesired effect the temperature of the cavity is held constant by a temperature control system.

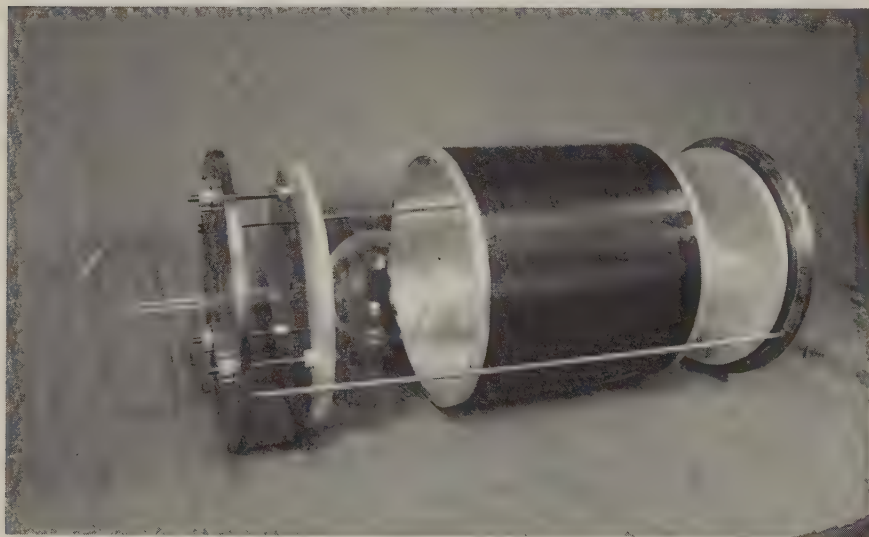


Fig. 37. Photograph of cavity parts.

The cavity is heated by a heater winding on the outer surface of the cavity cylinder. Enamelled copper wire is wound bifilarly since the heater current must not generate any magnetic field inside the cavity. For the same reason the wire diameter was chosen small (0.2 mm) to minimize

the current. A total wire length of app. 1400 m resulted in a resistance of about 750 ohms. Fig. 37 shows a photograph of the cavity taken apart where the heater winding can be seen.

The temperature of the cavity is measured by two thermistors mounted diagonally in the cavity cylinder (see fig. 38). The change in resistance is converted to a voltage in the input amplifier of the control system.

The temperature signal is fed to an analog input of a SATT Electronics PBS MIDI/C08 programmable process control system. A PID-controller generates an analog output signal to the voltage controlled power supply. The temperature of the cavity is held constant at approximately 41.5 °C by varying the power dissipation in the heater winding. To reduce the influence of room temperature fluctuations, the space between the cavity and the surrounding magnetic shields is filled by an insulating material.

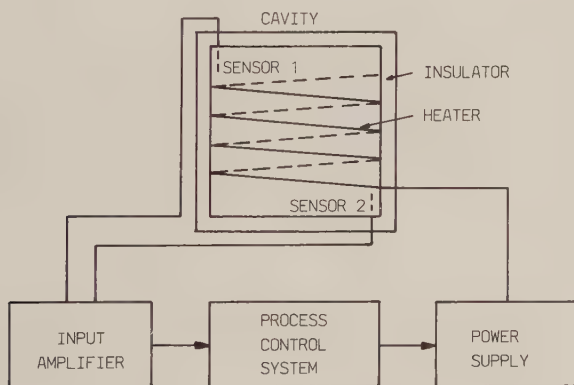


Fig. 38. Block diagram of temperature control system.

Preliminary tests show that it should be possible to keep the cavity temperature constant to within 0.001 degrees. This corresponds to a frequency deviation of 27 Hz which can be corrected by the cavity fine tuning servo. Further details and results will be published later by Göran Jönsson who is developing the temperature control system.

3.7 Magnetic field

An effect present in all atomic frequency standards is the magnetic field dependence of the transition frequency. The frequency sensitivity to changes in the magnetic field is higher for hydrogen than for other currently used atoms.

$$f = f_0 + 2750 \cdot 10^8 \cdot B^2 \quad (29)$$

where f_0 = frequency at zero magnetic field (Hz)
 B = average of the magnetic flux density
 inside the bulb (T)

It is therefore important for the stability of the standard that the magnetic flux density is held constant and uniform inside the bulb. These requirements are easiest to satisfy at very low magnetic fields, of the order of 10^{-7} T. At these low fields the influence of the earth's magnetic field (app. $5 \cdot 10^{-5}$ T) must be reduced as much as possible by use of magnetic shields.

3.7.1 Magnetic field coil

The most commonly used method of generating a homogenous magnetic field is to use a solenoid, i.e. a long cylindrical single layer coil of copper wire (fig 39). Other types of coils may be considered but are difficult to implement because of the limited space available inside the magnetic shields.

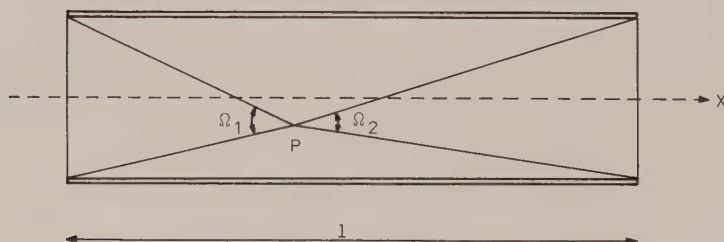


Fig. 39. Solenoid.

The component in the x direction of the magnetic flux density B_x at a point P inside a solenoid is <33>

$$B_x = \mu_0 \frac{N \cdot I}{l} \left(1 - \frac{\Omega_1 + \Omega_2}{4\pi} \right) \quad (30)$$

where I = coil current
 N = number of turns
 l = length of coil
 Ω_1, Ω_2 = space angles of the end surfaces

If the coil is long, Ω_1 and Ω_2 are small compared to 4π and B_x is approximately constant

$$B_x = \mu_0 \cdot \frac{N \cdot I}{l} \quad (31)$$

The actual length and diameter of the solenoid can, however, not be chosen freely. Both measures are restricted by the dimensions of the cavity and the magnetic shields. Fig. 40 shows the measurements of the manufactured solenoid.

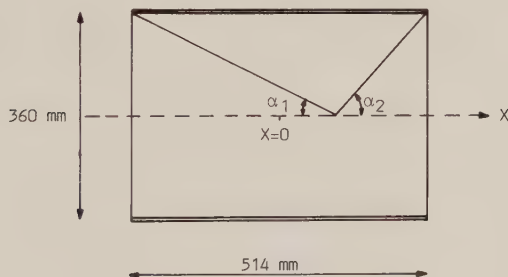


Fig. 40. Actual solenoid.

It is doubtful if the solenoid can be regarded as long compared to the diameter. The expression above for the magnetic field can be simplified to

$$B_x = \mu_0 \cdot \frac{N \cdot I}{l} \cdot \frac{1}{2} [\cos(\alpha_1) + \cos(\alpha_2)] \quad (32)$$

for points on the x-axis <33>. The diagram in fig. 41 shows the calculated relative magnetic field on the symmetry axis from a solenoid of the dimensions above with 790 turns of

0.65 mm enamelled copper wire. The plotted points represent the results of measurements made with an instrument using a Hall field probe.

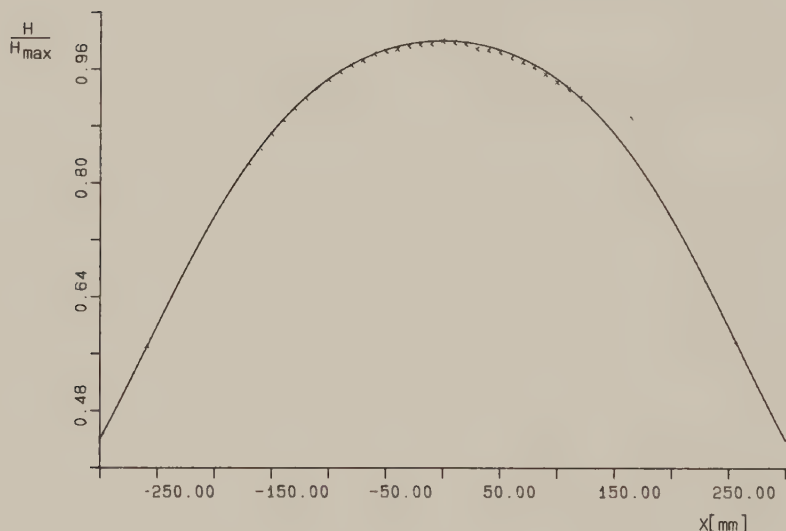


Fig. 41. Relative magnetic field on the symmetry axis from a solenoid.

It is obvious that the magnetic field is not constant enough in the area where the storage bulb is situated ($-70 < x < 70$).

There are two ways to compensate for the end effects. One way is to have taps on the coil so that the current in a number of turns at each end can be increased. The other way is to add compensation coils at the ends that are energized by the same current as the main coil. Since it seemed better to use only one current source, the second solution was chosen.

As an aid in designing the end compensation coils a computer program was written that calculated the sum of the magnetic flux densities on the x -axis from a solenoid and a selectable number of compensation coils. The number of turns for each pair of coils were optimized so that the magnetic field was constant over a maximum distance from the center point of the solenoid.

It was found that three pairs of compensation coils should be sufficient (fig. 42).

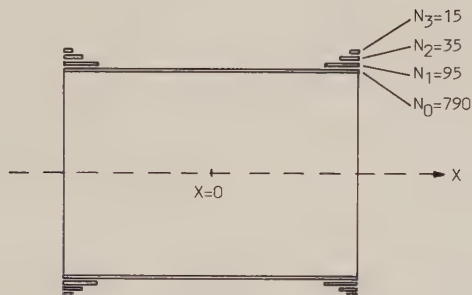


Fig. 42. Solenoid with compensation coils.

With 95, 35 and 15 turns on the coils the magnetic field could be held constant to within 0.1% in a ± 110 mm zone in the center. Outside this zone the field drops rapidly as the distance to the center increases. At ± 142 mm it has decreased 1% and at ± 215 mm it is 10% below the center value.

Measurements of the magnetic field on the center axis confirmed the calculations as shown in fig. 43. The improvement compared to the uncompensated solenoid is considerable.

The design of the compensation coils was based only on the magnetic field on the center axis because of the simple calculations. It was assumed, however, that the field away from the axis would improve in a similar way. This assumption was verified by a series of measurements along lines parallel to the center axis. The magnetic field inside a region corresponding to the volume of the cavity did not at any point deviate from the center value by more than 1%. The field orthogonal to the axis was less than 0.3%.

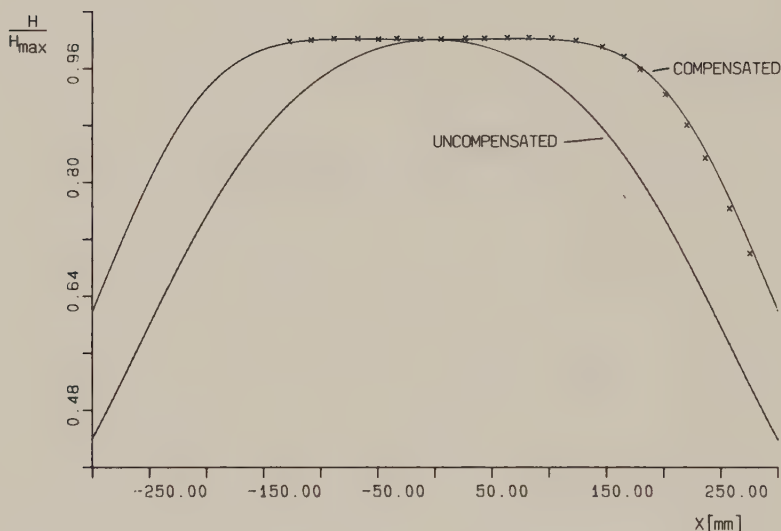


Fig. 43. Relative magnetic field on the symmetry axis from the compensated solenoid.

All the magnetic field measurements were performed in the presence of the earth's magnetic field. To reduce the influence of this constant field on the measurements, a high coil current was used so that the field from the coil was dominating. Furthermore, the mean value of two measurements, made with different current directions, was used to eliminate the constant field. The coil system was energized by a constant voltage from a standard stabilized laboratory power supply. These circumstances, in addition to the slow drift of the instrument used, limited the accuracy of the results to approximately 0.1%.

As seen from the equations above the magnetic flux density inside the compensated solenoid is directly proportional to the current. The proportionality factor could be calculated from the results of the measurements but would be of very little interest since the conditions are changed when the coil is put inside the magnetic shields. Furthermore the average magnetic field in the storage bulb can be measured more accurately using the magnetic field dependent Zeeman resonance frequency.

3.7.2 Constant current source

In order to maintain a constant magnetic field in the area where the interaction between the hydrogen atoms and the microwave field takes place, the solenoid must be fed with constant current. It is desirable to operate the maser at a low field strength but the residual magnetic field from external sources sets a lower limit. An operating field in the range $10^{-7} - 10^{-6}$ T is adequate if the shielding is sufficient. With the use of the compensated solenoid described in the previous section the required dc current is approximately 0.05-0.5 mA. Since the fine adjustment of the hydrogen maser frequency is done by changing the magnetic field, the current source must be continuously variable with sufficient resolution.

Since the output current is low and the voltage drop over the solenoid can be neglected due to the low coil resistance (78 ohms), a current source using operational amplifiers is suitable. The circuit in fig. 44 was chosen <34>.

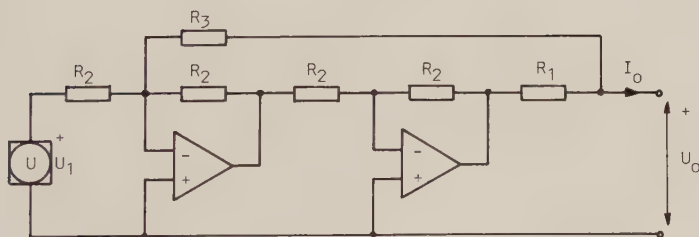


Fig. 44. Voltage controlled current source.

The output current I_O is given by the expression

$$I_O = \frac{U_1}{R_1} + \frac{R_2 - R_3 - R_1}{R_1 \cdot R_3} \cdot U_O \quad (33)$$

I_O is converted to a voltage by the resistor R_1 and is compared with the input voltage U_1 . If R_3 is made equal to $R_2 - R_1$ the output current is not dependent on the output voltage, i.e. the output impedance is equal to infinity. The output current is then determined only by the input voltage U_1 and the R_1 resistor.

$$I_O = \frac{U_1}{R_1} \quad (34)$$

If U_1 is a constant reference voltage the output current I_O is also held constant. The current can be changed by varying R_1 . To retain the high output impedance, however, R_3 must be adjusted simultaneously so that the condition $R_3 = R_2 - R_1$ is maintained. This is not a practical solution.

Another way to change the output current is to vary the input voltage U_1 . A third operational amplifier together with a precision voltage reference can be used to generate the variable input voltage (fig. 45).

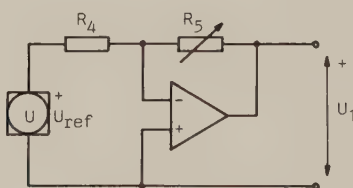


Fig. 45. Variable voltage reference.

The output voltage is changed by varying the ratio between R_5 and R_4 .

$$U_1 = - \frac{R_5}{R_4} \cdot U_{ref} \quad (35)$$

It is not advisable to vary R_4 since that would change the load resistance of the voltage reference, which could affect the output voltage. Therefore, a better solution is to vary R_5 . To accomplish both a wide setting range and a fine resolution, the resistance is varied in fixed steps by means of a switch and continuously by a multiturn preset potentiometer.

The voltage reference used is an AD580 2.5 volt high precision voltage reference with the following data:

Output voltage tolerance : < 75 mV
 Output voltage change : 85 ppm/°C
 Long term stability : 25 μ V/month
 Line regulation (V=23 V): < 6 mV
 Load regulation (I=10mA): < 10 mV

The following resistor values were used

$R_1 = 1 \text{ k}$
 $R_2 = 10 \text{ k}$
 $R_3 = 8.45\text{--}9.05 \text{ k} \quad (8.2 \text{ k} + 680 + 68 + 100(\text{pot}))$
 $R_4 = 2.13 \text{ k} \quad (1.8 \text{ k} + 330)$
 $R_5 = 0\text{--}10.02 \text{ k} \quad ((0 - 11) \cdot 820 + 1 \text{ k}(\text{pot}))$

which results in an output current range of 0–12 mA. The current can be changed continuously from 0 to 1 mA and up to 11 mA can be added in 1 mA steps by means of a switch. The 100 ohm potentiometer is adjusted so that the current is independent of the output voltage.

The fixed resistors are of the metal film type with an accuracy of 1% and a temperature coefficient of 50 ppm/°C. The potentiometers have 15 turns and a temperature coefficient of 100 ppm/°C. All amplifiers are LF351 JFET input operational amplifiers with low input offset voltage and very low input bias and offset currents.

The long term stability of the output current is mainly determined by changes in the reference voltage U_{ref} and changes in the resistance values of the resistors due to temperature variations. Stability can be increased if the temperature of these components is held constant. For this reason the complete current source is placed inside a small oven. The oven is heated to approximately 20 degrees above room temperature and held constant within 0.01 °C by a controller, also placed inside the oven. This arrangement should result in long term stability of less than $2.5 \cdot 10^{-6}$. This stability is difficult to verify directly from current measurements, and must be judged by the frequency stability of the frequency standard in operation.

3.7.3 Magnetic shield

Since it is desirable to operate the maser at a very weak magnetic field (10^{-7} T), it is necessary to reduce the contributions of other magnetic fields in the vicinity of the cavity to a level far below that generated by the solenoid. Examples of such fields are the earth's magnetic field ($0.5 \cdot 10^{-4} \text{ T}$) and the leakage fields from the permanent magnets of the ion vacuum pumps. The interfering fields can be attenuated by surrounding the cavity and the solenoid with a number of magnetic shields.

A magnetic shield is an enclosure made of a metal with a very high permeability μ . Because of the high μ the magnetic field is concentrated inside the metal wall and only a small fraction leaks to the inside (see fig. 46).

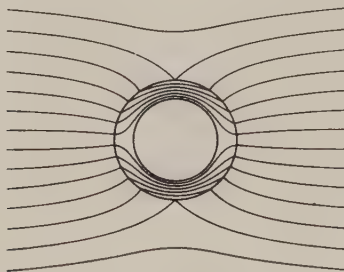


Fig. 46. Magnetic field streamlines for a cylindrical shield.

The shielding factor S is defined as the ratio between the outer field H_0 and the inner field H_i .

$$S = \frac{H_0}{H_i} \quad (36)$$

The shielding factor of a cylindrical shield can be calculated for a magnetic field perpendicular to the axis <35>.

$$S = \frac{\mu}{4} \left(1 - \frac{D_i^2}{D_o^2} \right) + 1 = \frac{\mu \cdot d}{D_o} \cdot \frac{D_m}{D_o} + 1 \approx \frac{\mu \cdot d}{D_o} \quad (37)$$

where μ = relative permeability

D_i = inner diameter

D_o = outer diameter

$D_m = \frac{D_i + D_o}{2}$ = mean diameter

$d = \frac{D_o - D_i}{2}$ = wall thickness

Even if wall thickness is increased the shielding factor will not exceed $S = \mu/4$ for a single shield. If a higher value is needed, two or more layers have to be used. The

resulting shielding factor for a double layer cylindrical shield is proportional to the product of the shielding factors for the individual single layers.

$$S = S_1 \cdot S_2 \cdot \frac{4\Delta}{D_o} \cdot \frac{D_m}{D_o} + S_1 + S_2 \approx S_1 \cdot S_2 \cdot \frac{4\Delta}{D_o} \quad (38)$$

where S_1 = shielding factor for shield 1
 S_2 = shielding factor for shield 2
 D_o = outer diameter
 D_m = mean diameter
 Δ = the distance between the shields

The shielding factor in a direction parallel to the axis of the cylinder is lower by a factor q .

$$S_{//} = q \cdot S \quad (39)$$

The reduction factor q depends on the ratio between the length L and the outer diameter D of the cylinder. If L/D is small q approaches unity and is gradually reduced for increasing L/D .

The shielding factors described above are only valid for static fields. The shielding factor for an ac field is in most cases higher due to the induced eddy currents. Furthermore, all results are valid only under the condition that the magnetic flux density is not so high that the μ -metal is saturated. If the shield approaches saturation the shielding factor decreases rapidly.

Since magnetic shields are very expensive and an old double layer shield from a previous frequency standard was available, it was decided to use this for the preliminary experiments. The two layers are made of μ -metal with a μ of 25000 and with the following dimensions:

Outer layer: $D = 450$ mm
 $L = 600$ mm
 $d = 2.0$ mm

Inner layer: $D = 402$ mm
 $L = 540$ mm
 $d = 0.8$ mm

Fig. 47 shows the shielding arrangement with the existing shields.

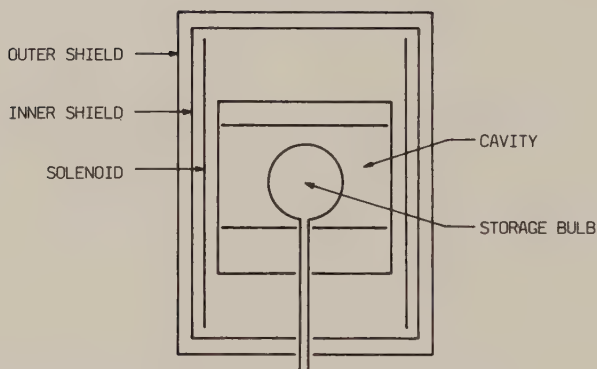


Fig. 47. Shielding arrangement.

The calculated shielding factor is 1170 for transverse magnetic fields and 840 for fields parallel to the axis. In practice it is necessary to have a number of openings in the end plates for the hydrogen beam, the microwave signals, mechanical supports to the cavity etc. These holes reduce the shielding factor depending on their size and design. Some additional holes had to be made for this new application. Tooling of the μ -metal destroys the high permeability, which has to be restored by annealing. This has not been done because of practical difficulties. Considering all these factors that decrease the shielding effect, the resulting shielding factor can be estimated to be of the order of 700 or 500 depending on the direction.

With an axial shielding factor of 500 and a vertical component of the earth's magnetic field of $0.5 \cdot 10^{-4}$ T, the residual field inside the cavity is approximately 10^{-7} T, i.e. of the same order as the desired operating field. If the magnetic field from the solenoid shall dominate over the background, the operating field must be at least 10 times higher, i.e. greater than 10^{-6} .

To be able to operate at the optimum magnetic field the shielding factor has to be improved considerably. This can be achieved only if the number of layers is increased to at least 4.

When used in a static field the shields will be permanently magnetized. The μ -metal can be demagnetized by applying a strong ac field that saturates the shields, and then slowly reducing the amplitude to zero. The ac field is induced by a three turn coil of thick copper wire running through the center holes of the end plates. The coil is energized by a current transformer able to deliver up to 500 A of 50 Hz alternating current.

3.8 Microwave signal generation

A 1420.405752 MHz microwave signal has to be derived from the 5 MHz crystal oscillator output frequency to enable the interaction with the hydrogen atoms. The desired frequency is generated by mixing a multiplied 1440 MHz signal ($288 \times 5 \text{ MHz}$) with a synthesized 19.594248 MHz signal (fig. 48).

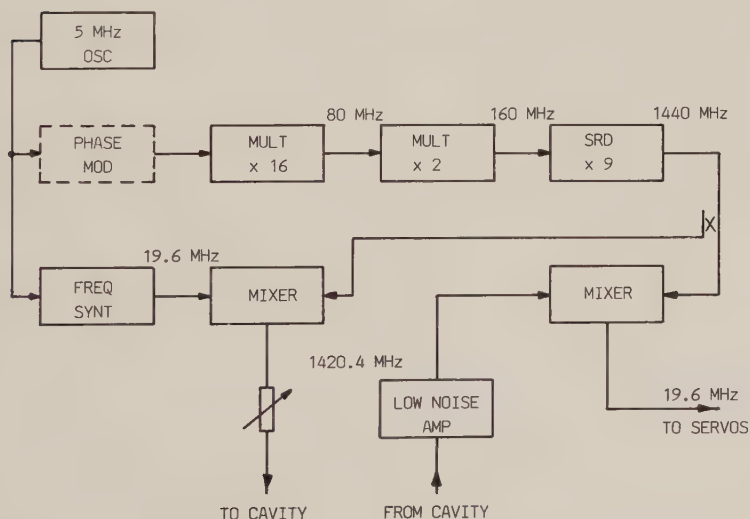


Fig. 48. Block diagram of microwave signal generator.

3.8.1 5 MHz crystal oscillator

An HP107BR quartz oscillator is used as the 5 MHz signal source. The quartz crystal and other critical components in the oscillator are housed in a temperature controlled double oven. The oscillator frequency can be electrically varied by means of a varactor diode.

The HP 107BR quartz oscillator has the following specifications of interest in this application:

Short term stability (1 sec)	:	$1.5 \cdot 10^{-11}$
Frequency control range (± 5 V)	:	$\pm 2 \cdot 10^{-8}$
Output voltage (into 50 ohms)	:	1 V _{rms}
Signal-to-noise ratio	:	> 87 dB
Harmonic distortion	:	< -40 dB

3.8.2 5 to 80 MHz multiplier

Several techniques can be considered for the design of frequency doubling stages at frequencies below 100 MHz. The two main principles are:

- Frequency multiplication using non-linear amplifiers with the output tuned to a multiple of the input frequency.
- Frequency doubling using double balanced mixers.

In this application, where the input frequency is phase modulated and the output signal must not have any amplitude modulation, circuits using high Q resonators are not suitable. Therefore the double balanced mixer seemed to be the best choice.

The diagram in fig. 49 shows a double balanced mixer based on the differential amplifier principle. Under the assumption that the transistors are ideal transconductance amplifiers, the currents in the load resistors can be calculated.

$$I_{C1} = \frac{I_E}{2} [1 - k_1 k_2 \sin(\omega_1 t) \cdot \sin(\omega_2 t)] \quad (40)$$

$$I_{C2} = \frac{I_E}{2} [1 + k_1 k_2 \sin(\omega_1 t) \cdot \sin(\omega_2 t)] \quad (41)$$

where $k_1 = g_m U_1$
 $k_2 = g_m U_2$
 g_m = transconductance of the transistors
 U_1, U_2 = input amplitudes
 ω_1, ω_2 = input frequencies

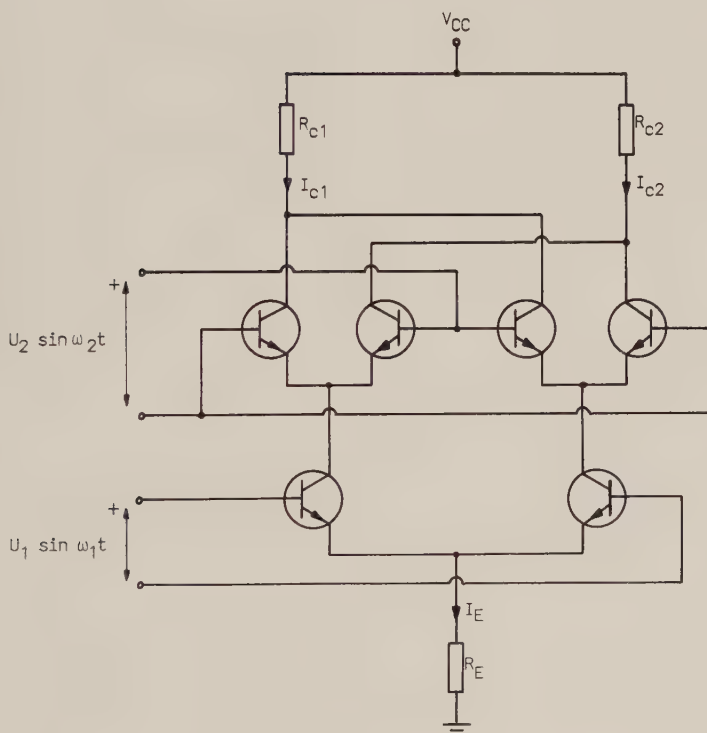


Fig. 49. Double balanced mixer.

If the input signals are identical the following simplified expressions are valid.

$$I_{C1} = \frac{I_E}{2} \left[1 - \frac{k_1 k_2}{2} + \frac{k_1 k_2}{2} \cos(2\omega t) \right] \quad (42)$$

$$I_{C2} = \frac{I_E}{2} \left[1 + \frac{k_1 k_2}{2} - \frac{k_1 k_2}{2} \cos(2\omega t) \right] \quad (43)$$

Except for the dc components, the load currents only contain components depending on 2ω , i.e. twice the input frequency.

As seen from the equations above there is no need to use tuned circuits in order to suppress components of either the input frequency or of harmonics of the input frequency. This is, however, only true if all the transistors are identical, which is not the case in reality. Therefore the input frequency will not be totally balanced out, and it might prove to be necessary to use low Q resonator loads to reduce the amplitudes of unwanted frequencies.

There are several types of integrated double balanced mixer circuits available from different manufacturers. A number of these circuits were tested and the S042 circuit from Siemens had the best behavior at high frequencies, good transistor matching, single supply operation and minimum requirements for external components.

The 5-80 MHz multiplier consists of an AGC input amplifier, four cascaded doubling stages and an output amplifier (fig. 50).



Fig. 50. Block diagram of 5-80 MHz multiplier.

The doubler modules using the S042 have the same basic design for all four stages, but with varying component values depending on the different operating frequencies (fig. 51).

The input signal is fed via capacitors C_3 and C_4 to one end of the two differential inputs. The other ends are grounded by capacitors C_5 and C_6 . A low Q resonator load consisting of R_1 , L_1 , C_1 and C_2 is used to suppress the input frequency.

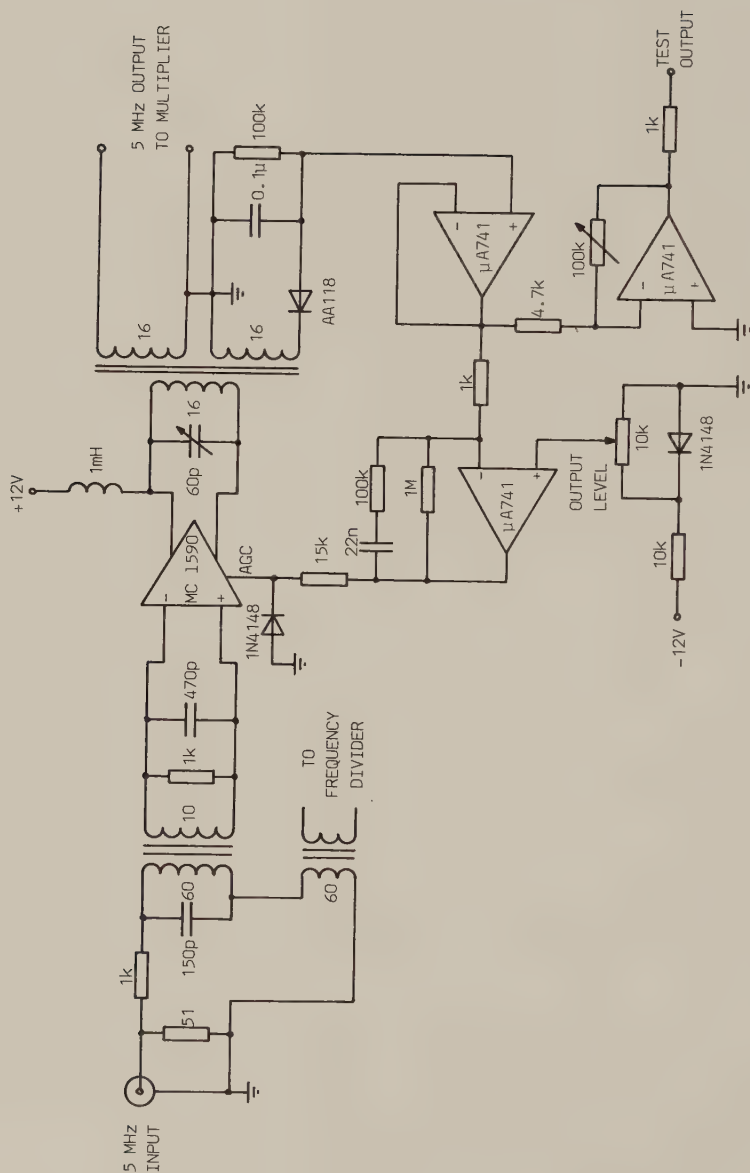


Fig. 52. 5 MHz AGC amplifier.

Special precautions concerning the printed circuit layout have been taken to prevent generation of undesired frequencies. The tuned circuits are placed far apart and separated by electromagnetic shields to reduce interstage crosstalk. The supply voltage is carefully decoupled at each doubler to avoid mixing via the supply voltage line.

It is essential to the performance of the multiplier that the signal amplitudes in the different stages are held constant at an optimum level. This is accomplished using an AGC amplifier on the 5 MHz input (fig. 52). Thus the input signal level to the first doubler stage is made independent of variations in the input signal amplitude.

The AGC amplifier consists of a transformer coupled MC1590 integrated differential amplifier. A third winding on the output transformer is used to detect the output amplitude. The output amplitude is controlled by a feedback loop from the amplitude detector to the gain control input of the amplifier. The output level is set by adjusting the reference voltage to the servo amplifier and can be monitored on the test output.

The output of the last doubler stage is buffered by an MC1590 rf amplifier in order to adjust the signal level to the following multiplier stages and to provide 50 ohm output impedance (fig. 53).

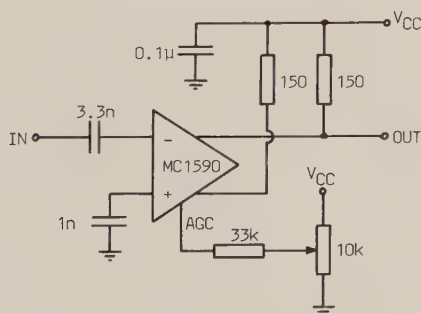


Fig. 53. 80 MHz buffer amplifier.

The AGC input voltage is used to manually control the gain by means of a potentiometer.

Spectrum analysis of the output signal gave the following results:

Output levels relative to the 80 MHz signal:

40 MHz	-34 dB
20 MHz	-39 dB
10 MHz	<-55 dB
5 MHz	<-55 dB

The Q value of the complete multiplier is 32 corresponding to a 3 dB output bandwidth of 2.5 MHz. If necessary the bandwidth can be increased at the expense of suppression of unwanted frequencies.

3.8.3 80 to 160 MHz frequency doubler

The currently used 80 to 160 MHz frequency doubler is an old construction designed for use in an earlier frequency standard <36>. The circuit diagrams are reproduced here, however, for completeness of the microwave signal generation description.

The frequency doubler can be divided into three functional blocks: a 80-160 MHz doubler stage using a tuned output transistor amplifier and two 160 MHz class C power amplifiers (fig. 54).



Fig. 54. Block diagram of 80 to 160 MHz frequency doubler.

The purpose of the power amplifiers (fig. 56 and 57) is to increase the signal amplitude to compensate for the relatively low efficiency of the succeeding step recovery diode multiplier.

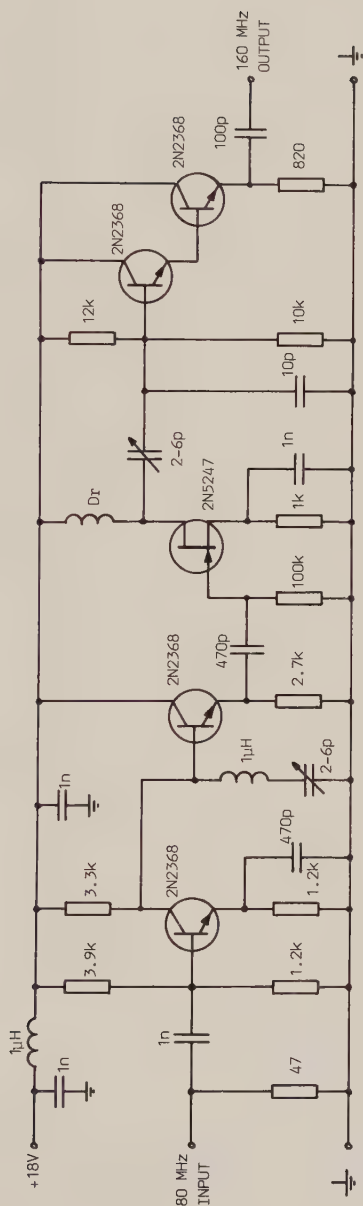


Fig. 55. 80 to 160 MHz frequency doubler.

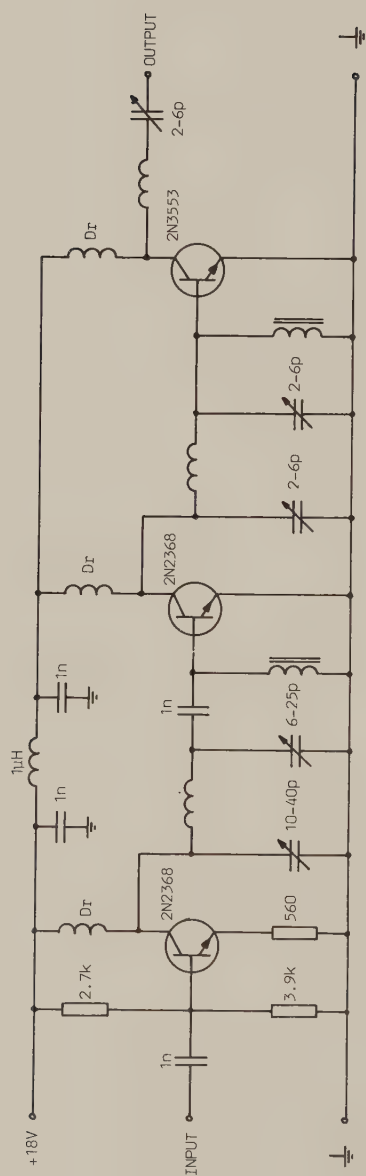


Fig. 56. First 160 MHz power amplifier.

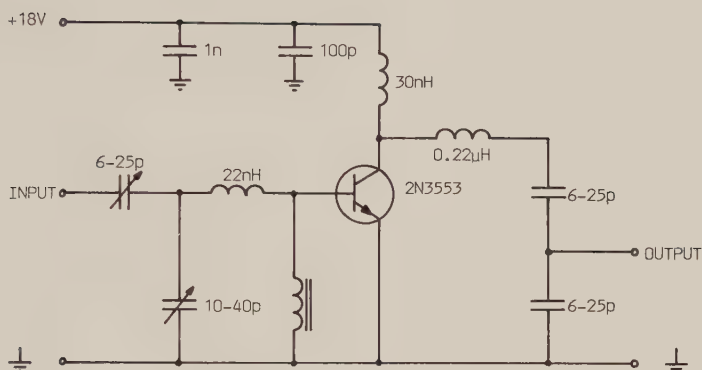


Fig. 57. Second 160 MHz power amplifier.

The 80 to 160 MHz frequency doubler delivers a 160 MHz output signal of 2 W with the frequency spectrum shown in fig. 58 and 59.

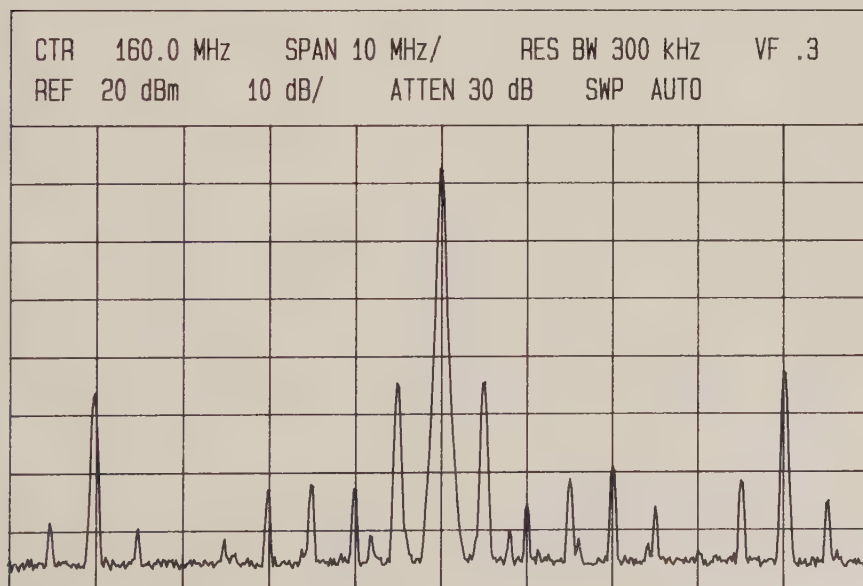


Fig. 58. 160 MHz output signal - 100 MHz span.

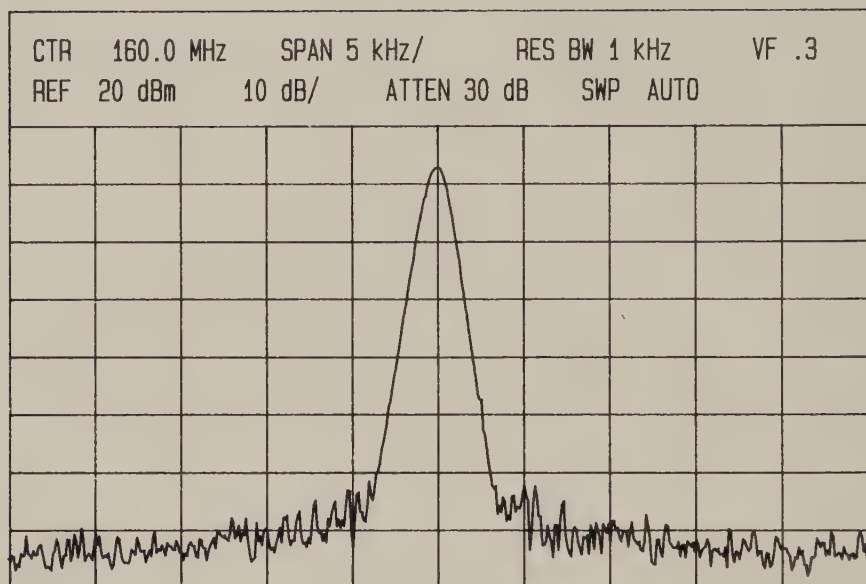


Fig. 59. 160 MHz output signal - 50 kHz span.

The spectra in fig. 58 and 59 are the results of measurements on the frequency doubler using a previous version of the 5-80 MHz multiplier as input signal. The old multiplier was designed using a different technique which gave a less clean output signal. The 40 MHz component was only 20 dB below the 80 MHz signal and there were -45 dB sidebands at ± 5 MHz from the carrier. This might explain the high peaks (app. -40 dB) at 120, 200, 155 and 165 MHz in the 160 MHz spectrum. The levels of these spurious frequencies should be lower using the new 5-80 MHz multiplier, but this has not been verified since I have been unable to obtain access to a spectrum analyzer since the new multiplier was implemented. The 3 dB bandwidth on the 160 MHz output is approximately 7.5 MHz.

3.8.4 160 to 1440 MHz SRD multiplier

The last multiplication by a factor of 9 from 160 MHz to 1440 MHz is done in a single stage using a step recovery diode <37>.

Energy at one frequency can be transferred to energy at a higher frequency utilizing the properties of an impulse. An ideal impulse has a frequency spectrum that contains all frequencies. If the impulse is repeated at a frequency f_1 , the frequency spectrum will have components at $n \cdot f_1$ where n is an integer. Generation of short pulses requires a very fast switch. The step recovery diode is such a switch.

The p-i-n doping profile (p-n junction separated by an intrinsic region) results in a highly nonlinear element. The diffusion capacitance in the forward direction is very large, resulting in a low impedance. In the reverse direction the diode capacitance is almost constant because of the intrinsic center region. As a simple first order model the SRD can be considered as a short circuit in the forward direction and as a capacitor in the reverse direction.

When the diode voltage is reversed the diode will continue to conduct a short time due to the charge stored during the forward bias period. When the charge equals zero the diode changes from the low impedance state to the high impedance state very abruptly (0.1 to 1 ns), which makes it possible to generate very short pulses.

The pulse generator circuit (fig. 60) consists of a voltage generator, a drive inductance, the step recovery diode, a bias battery and the load.

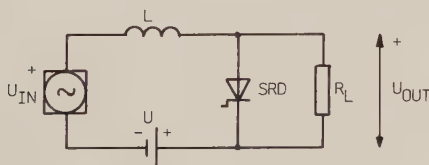


Fig. 60. Pulse generator circuit.

During the conduction interval the equivalent circuit in fig. 61 is valid.

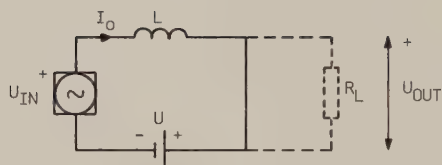


Fig. 61. Equivalent circuit during conduction interval.

The input current I_O flowing through the inductance L will vary with time as a function of the input voltage U_{in} and the battery voltage U . The output voltage is constant during the conducting phase and equal to the contact potential of the diode (0.7 V). When the charge stored during the forward bias period has been removed during the reverse bias period, the diode will cease to conduct. The time when this occurs is adjusted to the time when the current has its maximum negative value by varying the bias voltage U .

During the depletion interval another equivalent circuit is used (fig. 62), where the voltage sources are neglected.

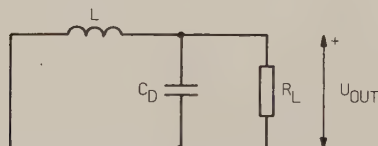


Fig. 62. Equivalent circuit during depletion interval.

If the capacitor were ideal, the current in the conductance would induce a damped sine wave signal across the output. The diode will be forward biased again after the first half period, however, and the ringing is short circuited by the diode. Thus a short pulse consisting of half a sine wave cycle appears on the output. The length of the pulse is determined by the diode capacitance C_D , the drive inductance L and the load resistor R_L . The height of the pulse is

limited by the breakdown voltage of the diode.

Fig. 63 shows the waveforms of the input voltage, the diode voltage and the input current.

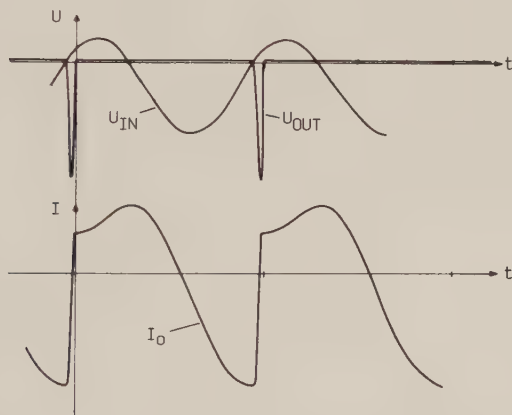


Fig. 63. Typical diode waveforms.

The HP 5082-0300 step recovery diode was found to be suitable for this application. It has a junction capacitance of 3.2-4.7 pF and a minimum breakdown voltage of 75 V.

For an output frequency of 1.44 GHz the drive inductance was calculated as 2.8 nH. The input impedance of the pulse generator consists of a resistance in parallel with an inductance. The inductance may be resonated at the input frequency by capacitor C_T , which also serves as a decoupling of the output frequency. A 180 pF capacitor makes the input impedance resistive and equal to 3.1 ohms. The low input impedance is matched to the 50 ohm output impedance of the 160 MHz power amplifier by an L-C network consisting of L_1 , C_1 and C_2 (fig. 64).

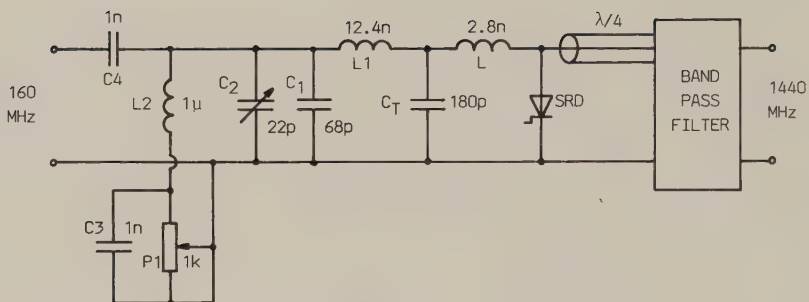


Fig. 64. Diagram of the complete SRD multiplier.

The input section is completed by the circuit L_2 , C_3 and P_1 for automatic generation of the bias voltage and the dc blocking capacitor C_4 .

The output energy may be concentrated to the desired output frequency by converting the short pulse to a damped sine wave signal. This conversion is accomplished by feeding the pulse into a short coaxial transmission line. If the transmission line is terminated by a resistance larger than its characteristic impedance, the pulse will be reflected back towards the diode. If the length of the line is a quarter of a wavelength the diode will be conducting again when the pulse arrives. The diode acts as a short circuit and reflects the pulse once more with an inverted phase. Thus the pulse will travel back and forth forming a damped sine wave at the termination. If the load resistance is adjusted so that the ringing is just damped out when the next pulse is generated by the diode, all the pulse energy is delivered to the load, but with a changed spectral distribution.

The damped sine wave may be filtered by a coaxial cavity bandpass filter to obtain a constant output amplitude. Fig. 65 shows the mechanical design of the complete step recovery diode multiplier <37>.

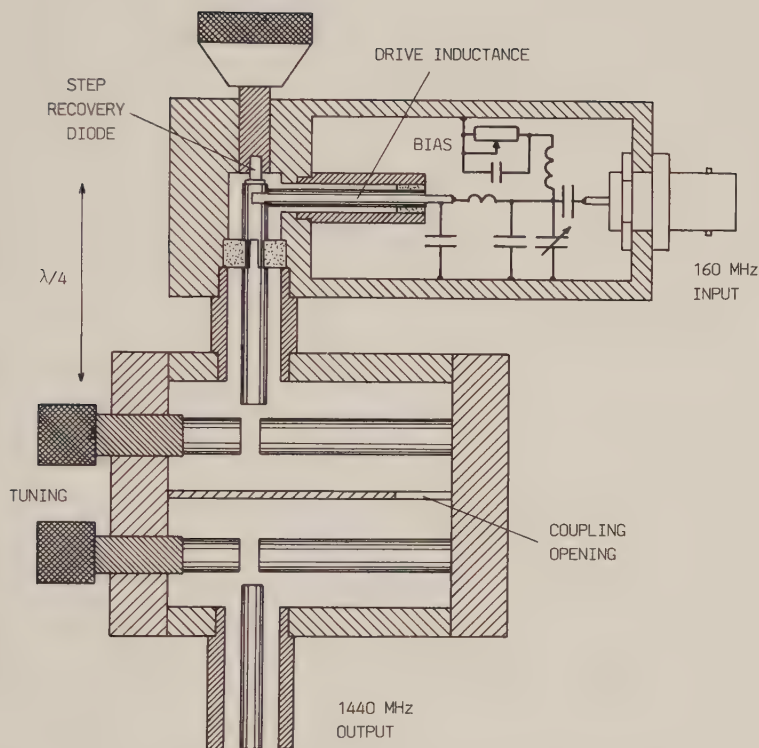


Fig. 65. Mechanical design of SRD multiplier.

The diode is mounted at the end of the quarter wave transmission line. All input circuit components are discrete except for the drive inductance which is implemented as a short length of 32 ohm transmission line. The input and output of the bandpass filter are electrically coupled to the coaxial resonators by the extended inner lead of the 50 ohm transmission lines. The resonators are coupled by means of an opening in the partition wall. The cavities are mechanically tuned by varying the gap between the center posts.

The shape of the pulses (see fig. 66) were recorded using a sampling oscilloscope with the bandpass filter section removed and replaced by a 50 ohm termination. The amplitude of the pulses is 35 V and the width at half the amplitude is 0.54 ns. The damped sine wave signal was measured using the quarter wave line and a loose coupling to the oscilloscope. Fig. 67 shows the 9 periods of 1440 MHz signal with decreasing amplitude in the period between the 160 MHz pulses. Due to the loose coupling the ringing is not completely damped out when the new excitation occurs.

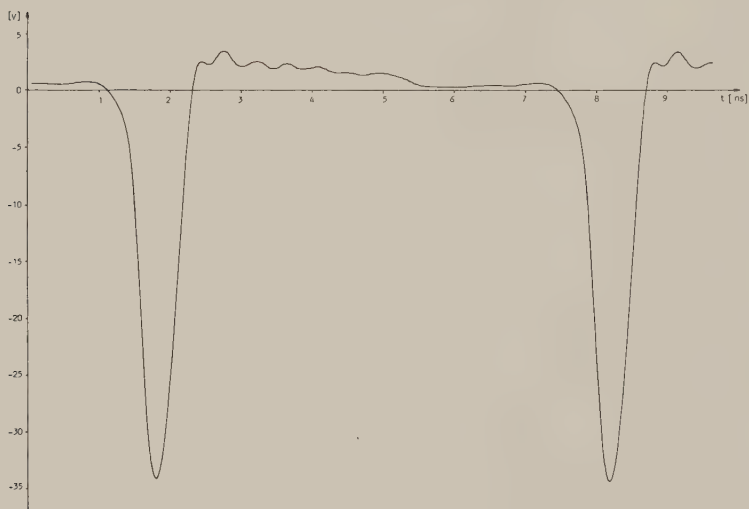


Fig. 66. Output signal from the SRD pulse generator circuit.

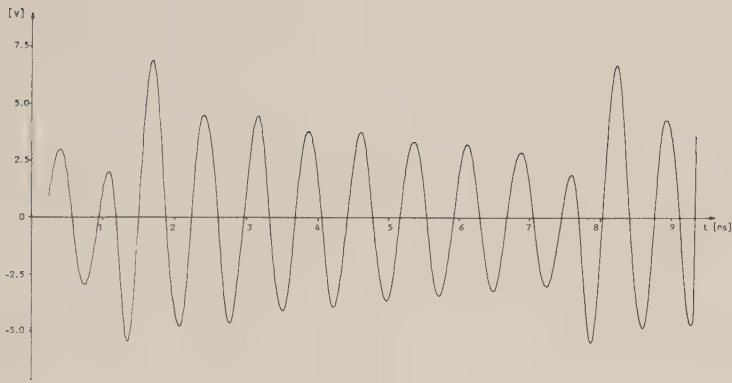


Fig. 67. Output signal from the quarter wave line.

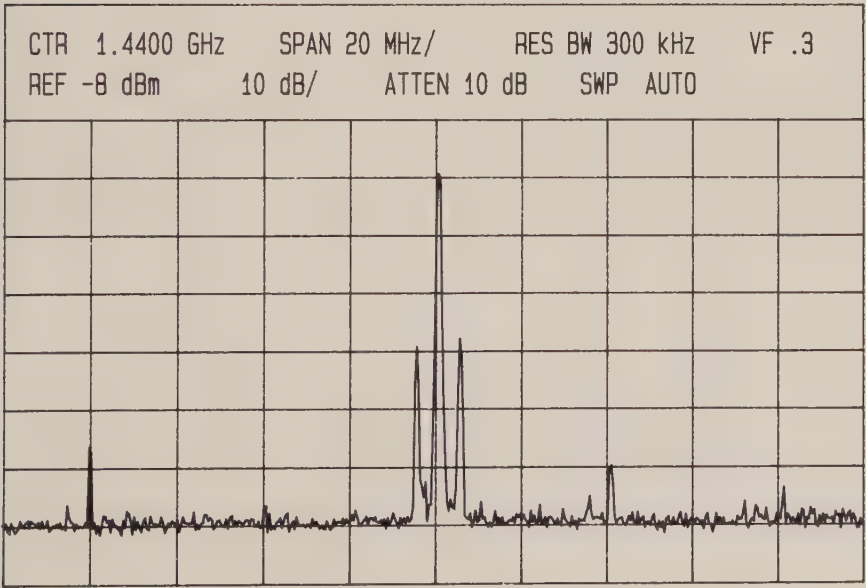


Fig. 68. Output frequency spectrum from the SRD multiplier - 200 MHz span.

The frequency spectrum of the output signal (fig. 68) shows that the sidebands 5 MHz away from the carrier are still present at a level 30 dB below the carrier. These sidebands are probably originating from the old 5-80 MHz multiplier used when these measurements were made and should be smaller using the new version. They are too close to the carrier to be filtered by the bandpass filter which has a 3 dB bandwidth of approximately 4 MHz.

Fig. 69 and 70 show the frequency spectrum close to the carrier with and without the 12.5 kHz modulation for the cavity control.

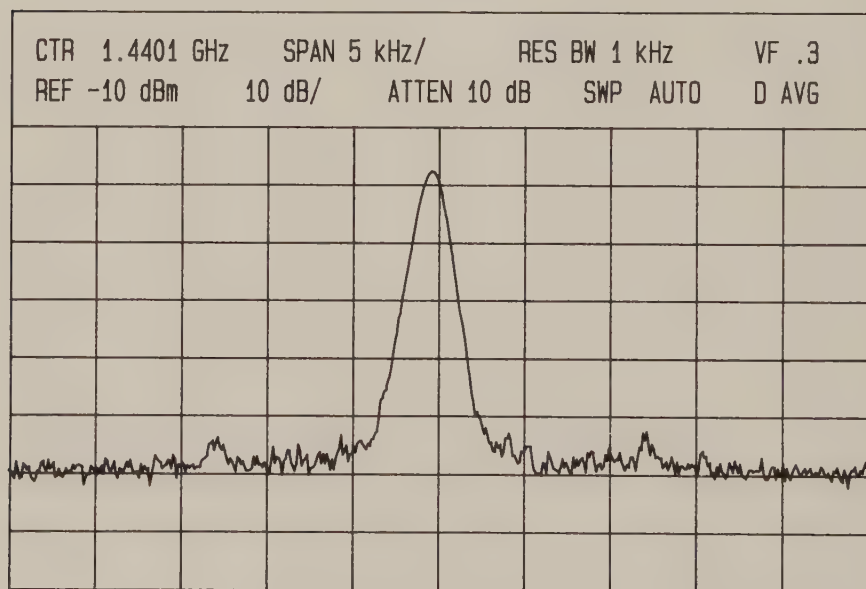


Fig. 69. Output frequency spectrum from the SRD multiplier - 50 kHz span - without 12.5 kHz modulation.

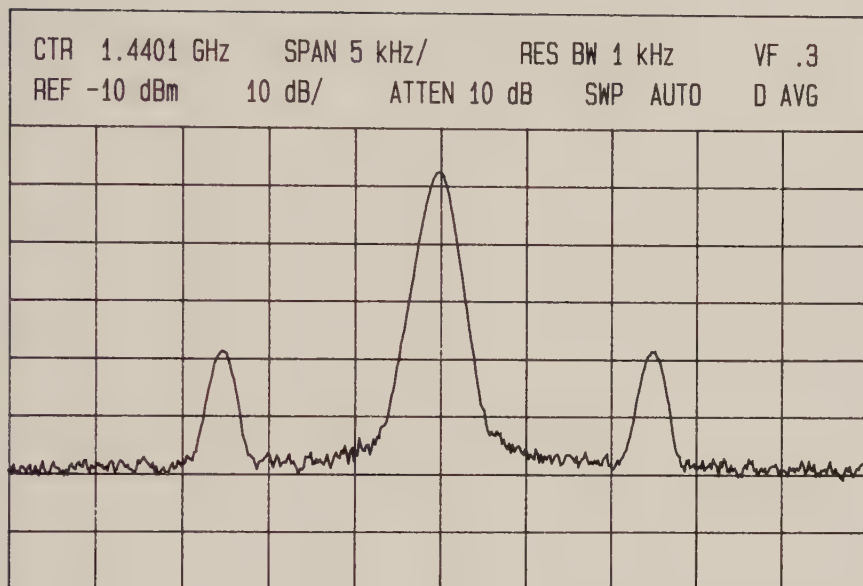


Fig. 70. Output frequency spectrum from the SRD multiplier - 50 kHz span - with 12.5 kHz modulation.

The efficiency of the step recovery diode can according to the manufacturer be empirically estimated to

$$\eta = 100 \cdot \frac{K}{N} \quad (44)$$

where $K = 1-2$ depending on the output frequency
 $N =$ the harmonic number

For a multiplication factor of 9 the efficiency should be in the range of 11 to 22 percent. The output power of the complete SRD multiplier was measured to 400 mW (+26 dBm) at an input power of 2 W. This is equivalent to an efficiency of 12.5 % which is satisfactory. The highest 1440 MHz signal level, +7 dBm, is required at the local oscillator input of the down conversion mixer.

3.8.5 Mixing to the hydrogen line frequency

The exact frequency of the hydrogen line is generated as the difference frequency between the 1440 MHz signal from the multiplier chain and a synthesized signal of 19.6 MHz (fig. 71).

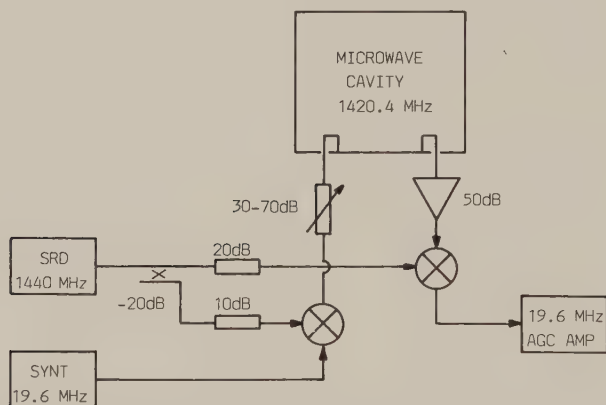


Fig. 71. Schematic diagram of rf section.

The output signal from the SRD multiplier is divided into two paths, one for generation of the cavity input signal and one for down conversion of the output signal from the cavity. The highest signal level is required by the local oscillator port of the down conversion mixer (+ 7 dBm) and only a small part of the available signal power is needed for the input frequency generation. Therefore the latter signal is tapped from the local oscillator signal by a directional coupler with an output level of -20 dB compared to the main output. The signals are adjusted to suitable levels by fixed coaxial attenuators, except for the input signal to the cavity, which may be varied continuously.

Doubly balanced mixers (DBM-400A from Vari-L) with the following specifications are used for the frequency conversions.

Data for DBM-400A:

Frequency range	RF port	100-3000 MHz
	LO port	10-3000 MHz
	IF port	DC-1000 MHz
Isolation	LO-RF	> 35 dB
	LO-IF	> 30 dB
	RF-IF	> 20 dB
Recommended LO power		+ 7 dBm
Conversion loss		< 8.5 dB

Fig. 72 shows the output frequency spectrum from the mixer generating the input frequency to the cavity. The unwanted frequencies are not a problem due to the narrow bandwidth of the cavity.

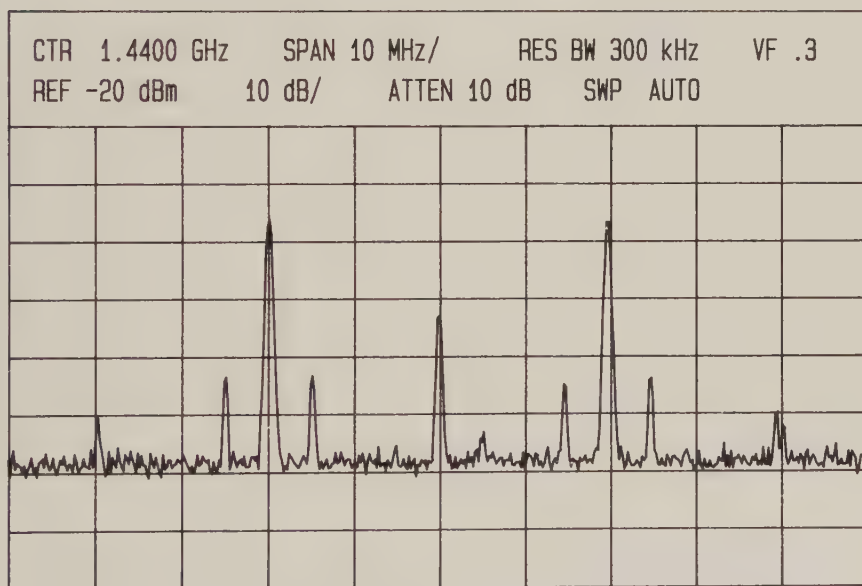


Fig. 72. Frequency spectrum of the cavity input signal.

Because of the very low output power from the microwave cavity (on the order of -100 dBm) the signal must be amplified in a low noise preamplifier before it is fed to the down conversion mixer. An AMG-2022N wideband low noise figure amplifier from Avantek is used for this purpose. The following data were measured by the manufacturer at a frequency of 1500 MHz:

Gain 50.6 dB
 Noise figure 2.55 dB
 VSWR < 2.0

Fig. 73 shows the 1420.4 MHz output signal after amplification by the low noise preamplifier.

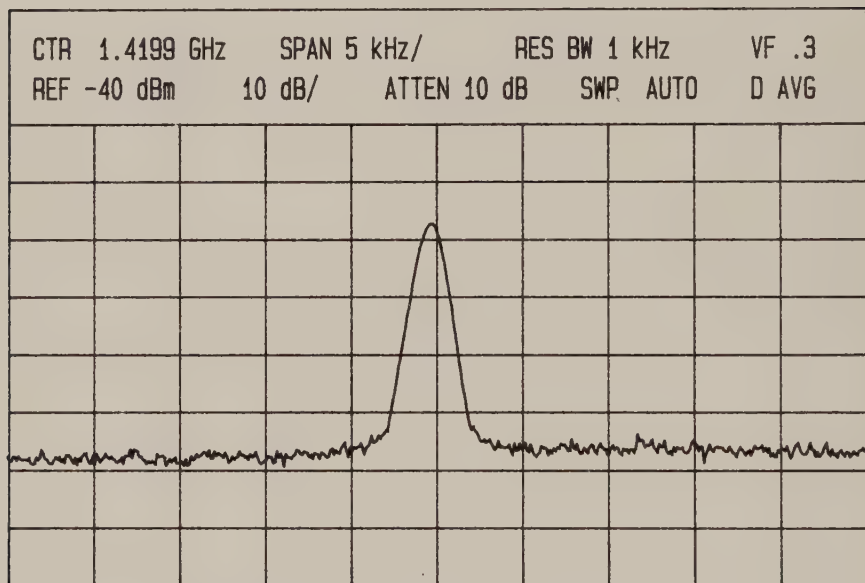


Fig. 73. Frequency spectrum of the low noise amplifier output signal.

The synthesized 19594248 Hz signal to be mixed with the 1440 MHz signal from the multiplier is supplied by an HP 3335A frequency synthesizer using the 5 MHz oscillator signal as an external frequency reference. It was originally intended to build a special synthesizer for this purpose, and several simple prototypes have been built to evaluate different designs. This work was never completed, however, since other parts of the standard were considered to be of higher priority.

3.9 Cavity and hydrogen line servos

The control electronics of the passive hydrogen maser consist of two frequency locked loops:

- The cavity control loop for tuning the cavity to exactly the same frequency as the hydrogen line to reduce the cavity pulling effect.
- The hydrogen line servo loop for locking the 5 MHz crystal oscillator to the hydrogen line to increase the long-term stability.

Both loops use phase modulation of the microwave signal to probe the cavity and hydrogen resonances. The frequency dependence of the output amplitude due to the resonances results in an amplitude modulation in addition to the phase modulation. The phase of the amplitude modulation compared to the phase of the modulating signal gives information about on which side of the resonance the center frequency of the sweep is located. This information is used to align the center frequency of the sweep and the center frequency of the resonance.

The cavity and hydrogen line servos consist of the following functional blocks (fig. 74).

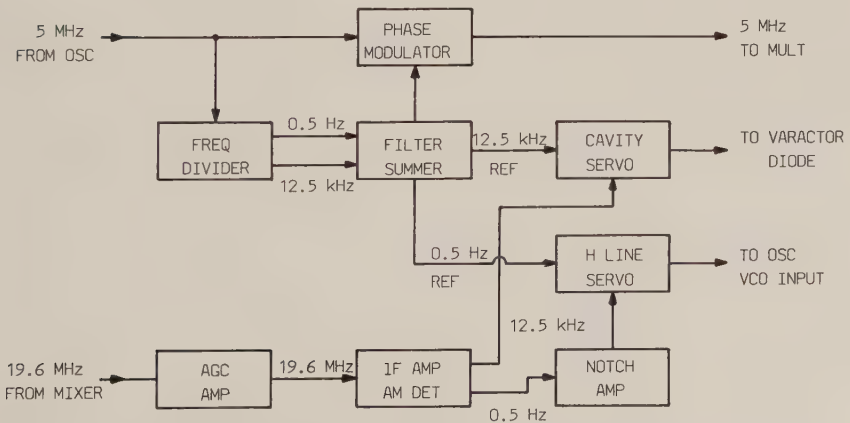


Fig. 74. Block diagram of cavity and hydrogen servos.

Two separate modulation frequencies are used: 0.5 Hz for the narrow hydrogen resonance and 12.5 kHz for the wider cavity

resonance. Both modulation frequencies are derived from the 5 MHz oscillator signal by means of a digital frequency divider. The two frequencies are converted to sine waves in bandpass filters and added in appropriate portions to one modulation signal. The modulation signal is fed to the phase modulator which is modulating the 5 MHz signal before it is multiplied to 1440 MHz and mixed with the synthesized signal to the hydrogen transition frequency. The output signal from the cavity is converted down to 19.6 MHz before it is fed to the tuned AGC amplifier in the receiver section of the servos. In the if amplifier the signal is filtered and divided into two signal paths, one for each modulation frequency. The two modulation frequencies are detected in separate amplitude detectors, amplified and compared to the corresponding reference signals in the phase sensitive detectors. The phase information is converted by servo amplifiers to control voltages for the varactor diodes in the cavity tuning and the 5 MHz crystal oscillator, thus closing the servo loops.

The hydrogen line and cavity servos are only briefly described by block diagrams illustrating the basic operating principles and the design goals. The details of these parts are designed by Bo Olsson (except for the modulation frequency generator) and are described in detail in a separate report <41>.

3.9.1 Modulation frequency generation

Both phase modulation frequencies, 0.5 Hz and 12.5 kHz, are derived from the 5 MHz oscillator signal by means of a digital frequency divider. The sine wave input signal must be converted to a square wave with TTL levels before it can be used as a clock signal for the counters. This conversion is made using an NE527 high speed analog voltage comparator with TTL compatible outputs (fig. 75). To minimize changes in phase and duty cycle due to variations in the sine wave amplitude, the comparator is preceded by an AGC amplifier which keeps the input amplitude constant. This AGC amplifier is identical to the input amplifier of the 5-80 MHz multiplier and uses the same 5 MHz input. The input of the comparator is connected to the amplifier output via a differential transformer to eliminate offset problems. The output of the comparator is buffered by two 50 ohm line drivers.

The frequency divider (fig. 76) consists mainly of six synchronous decade counters. The last steps in the modulation frequency generation are divide-by-two circuits using edge triggered J-K flip-flops. This ensures that the output waveforms are totally symmetric and contain no second harmonic. The differential outputs are buffered by 50 ohm line drivers.

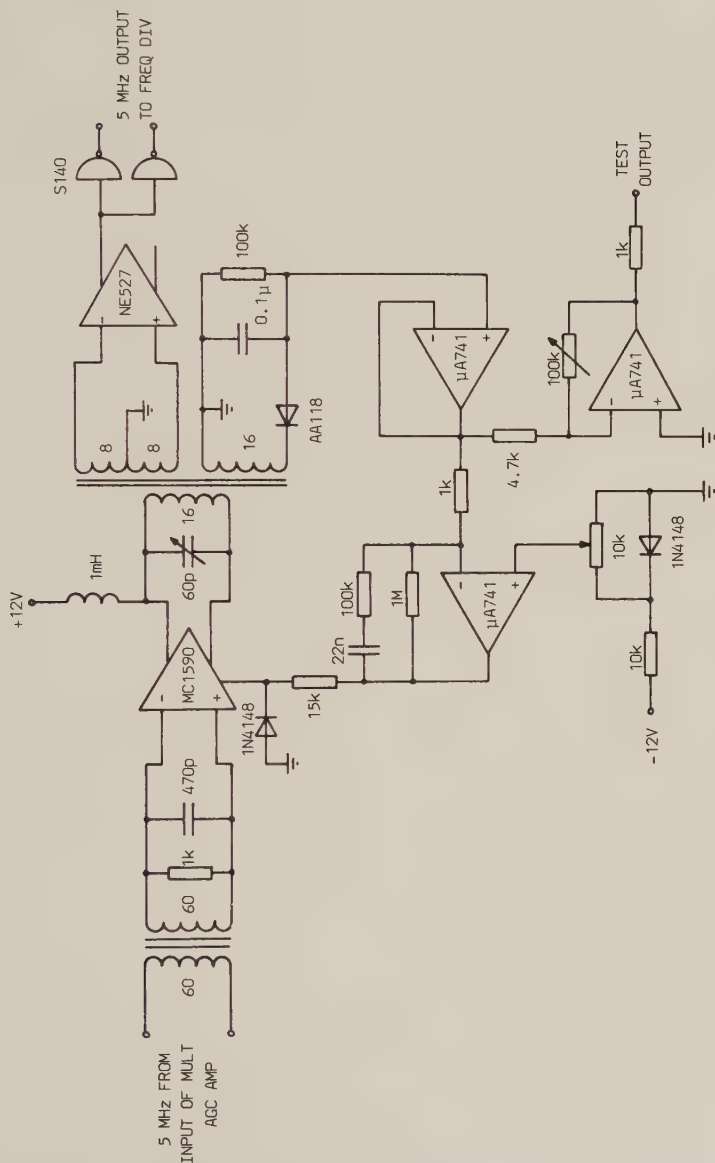


Fig. 75. 5 MHz sine wave to square wave converter.

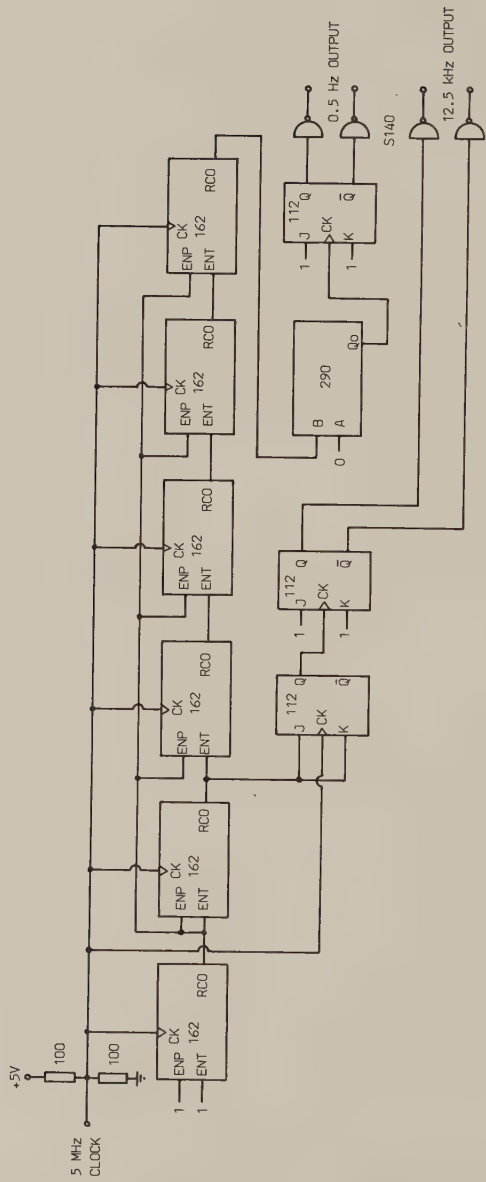


Fig. 76. Modulation frequency generator.

3.9.2 Phase modulator

The 0.5 Hz and 12.5 kHz square wave signals from the frequency divider are converted to sine wave signals using active bandpass filters (fig. 77).

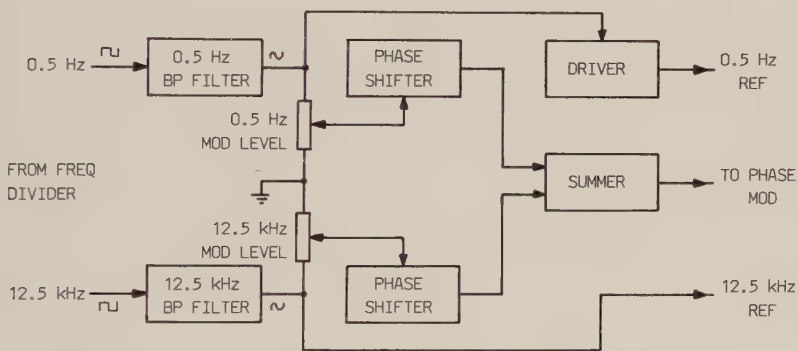


Fig. 77. Filter and summing circuit.

The amplitudes are adjusted according to the desired degree of modulation, phase shifted 90 degrees and combined into one modulating signal in a summing amplifier. This section also has outputs for 0.5 Hz and 12.5 kHz reference signals for the synchronous detectors.

The phase modulator section consists of a phase modulator and an output amplifier (fig.78).



Fig. 78. Phase modulator.

The phase modulator is an allpass network using varactor diodes as a variable capacitor <42>. The phase modulation angles are small, specially for the 0.5 Hz modulation, which makes it possible to satisfy the very high demands for suppression of amplitude modulation. It is of utmost importance for the performance of the frequency standard

that the amplitude modulation of the output signal from the cavity originates from the hydrogen line and not from the phase modulator itself. The output amplifier is required to restore the original signal amplitude and to drive a 50 ohm load.

3.9.3 19.6 MHz AGC amplifier

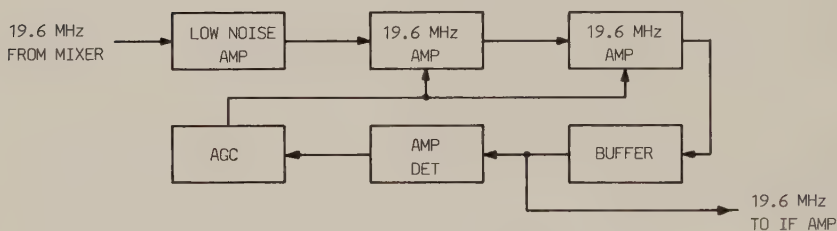


Fig. 79. 19.6 MHz AGC amplifier.

The 19.6 MHz input signal from the down conversion mixer is amplified to a constant output level (fig. 79). The gains of the second and third tuned amplifier stages are controlled by an automatic gain control loop.

3.9.4 Amplitude detectors

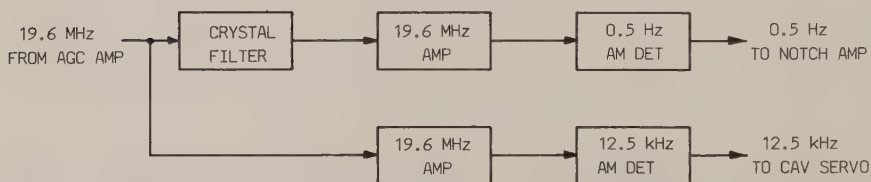


Fig. 80. Amplitude detectors.

From this point the signal paths for the 0.5 Hz and the 12.5

kHz modulation signals are separated (see fig. 80). The 0.5 Hz signal path contains a narrow bandwidth crystal filter to remove the 12.5 kHz sidebands. After an additional amplifier the 0.5 Hz amplitude modulation is detected in a conventional diode detector. The 12.5 kHz signal path has a similar design except for the crystal filter.

3.9.5 Hydrogen line servo

The 0.5 Hz signal is amplified further in a notch amplifier (fig. 81) suppressing the 1.0 Hz second harmonic generated by the hydrogen line. The total gain of the amplifier is selectable between 1000 and 50000.

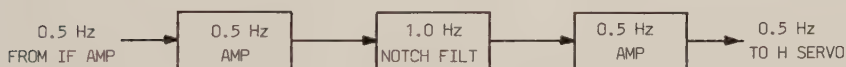


Fig. 81. Notch amplifier.

The phase of the detected 0.5 Hz amplitude modulation is compared to the phase of the reference signal in a synchronous detector (fig. 82). The phase error signal is fed to the servo amplifiers which produce a control signal to adjust the 5 MHz frequency until the phase error is eliminated.

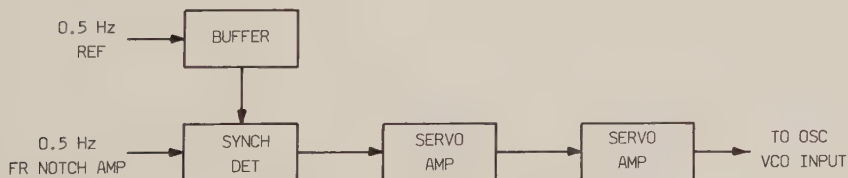


Fig. 82. Hydrogen line servo.

The time constant of the servo is chosen in such a way that the short-term stability of the crystal oscillator and the long-term stability of the hydrogen line are combined.

3.9.6 Cavity servo

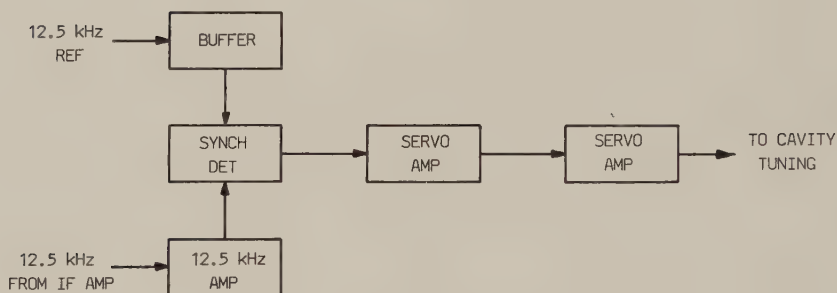


Fig. 83. Cavity servo

The structure of the cavity servo (fig. 83) is similar to that of the hydrogen line servo. The 12.5 kHz input amplifier is of a less complex design and the time constant of the servo is smaller due to the higher modulation frequency. The output voltage adjusts the resonance frequency of the cavity until it coincides with the hydrogen line frequency.

4 FREQUENCY DIVIDER IC FOR TICK GENERATION

4.1 Introduction

In the field of time and frequency it is often of interest to generate a one pulse per second tick signal from the 5 MHz sine wave output from an atomic frequency standard. The tick signal can be used for time keeping purposes or in long-term stability comparisons between different standards by time interval measurements.

Existing designs for tick generators using standard TTL or CMOS counter circuits show certain drawbacks that could be eliminated using modern large scale integration technology. A fully integrated tick generator would feature higher reliability, lower power consumption, smaller size etc. On these grounds it was decided that an attempt should be made to design such an integrated circuit. The available SOS/CMOS process was well suited for this high speed digital application.

The following design goals were established:

- The integrated CMOS divider circuit should feature higher reliability and less sensitivity to disturbances than previous tick generators.
- The input signal should be a 5 MHz sine wave with varying amplitude taken directly from a frequency standard.
- The output signal should be 1 PPS pulses with short rise- and fall-times and very low jitter.
- The circuit should have a reset input so that two or more tick generators could be synchronized.
- The circuit should also produce symmetrical square wave output signals for all decades between 1 MHz and 1 Hz.
- The circuit should, if possible, also be able to accept a 10 MHz input signal.
- The circuit should have small dimensions to allow a number of circuits to be coupled in parallel to achieve further increased reliability.

4.2 Functional description

The basic functions of the circuit are described by the block diagram in fig. 84.

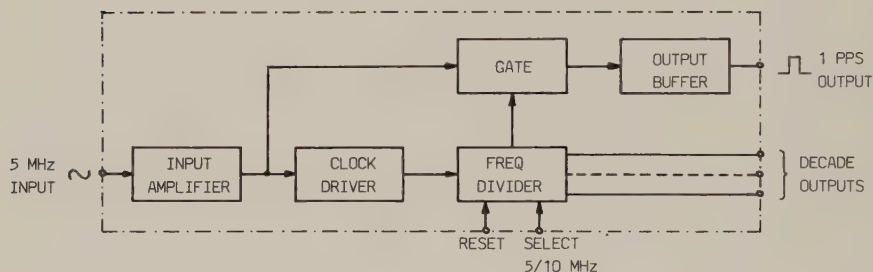


Fig. 84. Block diagram of tick generator circuit.

The input 5 (or 10) MHz sine wave signal is converted into a square wave signal by the input amplifier. The gain of this amplifier must be high enough to ensure that the slopes of the square wave are steep even if the amplitude of the input signal is fairly low. The input frequency is fed to a divider circuit consisting of synchronous decade counters that generate a short pulse once every second. This pulse gates one period of the input signal to the output. Thus the delay and ripple of the counters do not affect the output pulse. As the clock signal must drive a large number of gates at a high frequency it is necessary to insert a clock driver between the input amplifier and the counters. Also the output pulse should be buffered to be able to drive a large external capacitive load. If the decade counters are designed to generate symmetrical divide-by-ten output signals, it is possible also to get square wave outputs for 1 MHz, 100 kHz, etc. down to 1 Hz.

The frequency divider (fig. 85) which divides the input frequency by a factor 5 000 000 (10 000 000 if selected) contains 7 synchronous, resettable 4 bit counters.

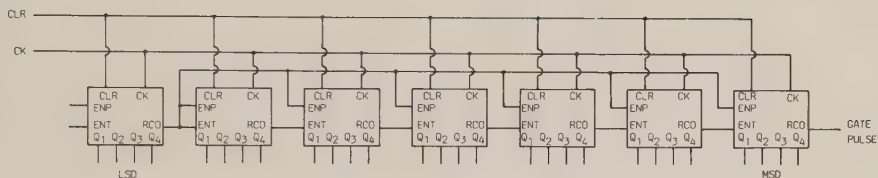


Fig. 85. Frequency divider.

The counters feature an internal carry look-ahead to enable cascading at high frequencies. Synchronous operation is provided by clocking all counters simultaneously so that all the outputs change in unison. The clear input resets the counters synchronously at the next clock pulse. The ripple carry output from the most significant decade is used as the gating pulse. The carry pulse is high when the contents of the counters are equal to 4 999 999 (or 9 999 999).

The decade counters (fig. 86) consist of 4 flip-flops and some logic gates to generate the desired counting sequence and to handle the carry signals. Since all basic gates in CMOS have inverting outputs, the logic circuitry is designed using exclusively NAND-gates with varying numbers of inputs. NAND-gates also have lower output impedance compared to NOR-gates which reduces the gate delay.

Since no NBCD output of the contents of the decade counter is required, but merely a division of the incoming frequency by ten, the counting sequence (fig. 87) is changed a bit to provide a symmetrical output signal. The upper three flip-flops divide by 5 and the remaining flip-flop divides by 2.

The counter counts only when both the enable signals (ENP and ENT) are high and the clear signal (CLR) is low. The ripple carry output (RCO) goes high when the content of the counter equals nine (Q_3 and Q_4 high) and ENT is high.

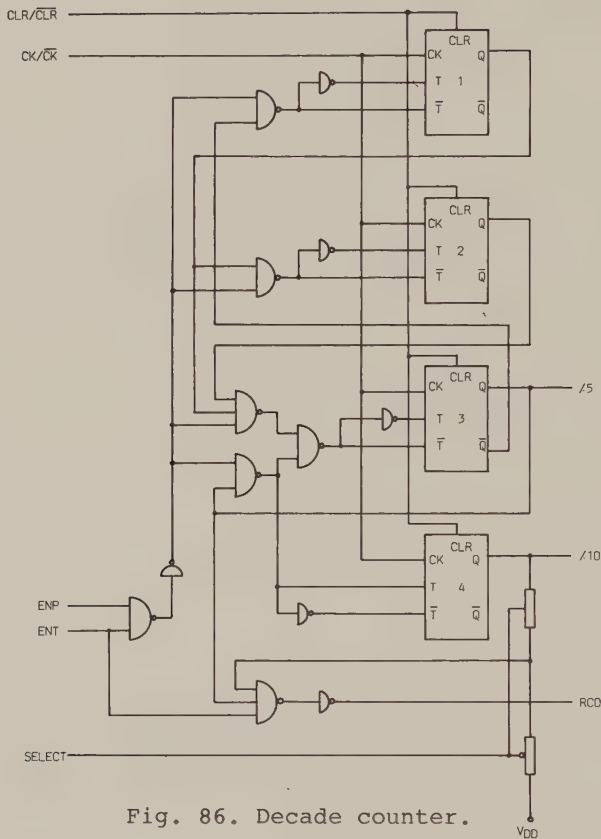


Fig. 86. Decade counter.

As an additional feature the most significant flip-flop can be electrically 'disconnected' by two transmission gates controlled by the select (SEL) signal. If SEL is low the counter runs in divide-by-five mode and the carry output goes high when the counter reaches 4 (Q_3 high).

The counter uses a special carry arrangement where the carry chain is divided into two separate paths. One path passes through all the decade counters from the least significant to the most significant. In each of the decades the signal is delayed a time proportional to the number of gates through which it must propagate. The total delay of the carry signal limits the maximum counting speed considerably as the number of decades is increased. This effect, however, is eliminated by the second path, where the carry output from the least significant decade is connected directly to the ENP inputs of all the other decades in parallel. Thus the effective carry signal delay is small and

independent of the number of decades.

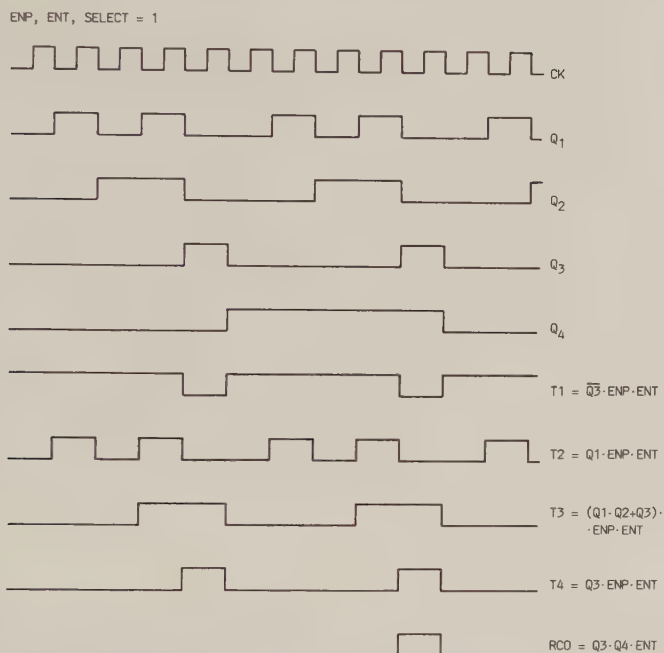


Fig. 87. Counting sequence.

Further details on the design of the decade counters can be found in the appendix, which also contains the results of computer simulations and breadboarding of the basic circuits and a detailed description of the layout generation.

4.3 Performance evaluation

The tick generator circuit was processed as a part of a multiproject chip (MPC) containing a large number of layouts of varying size and complexity. Fig. 88 and 89 show photographs of the processed frequency divider chip taken through a microscope.

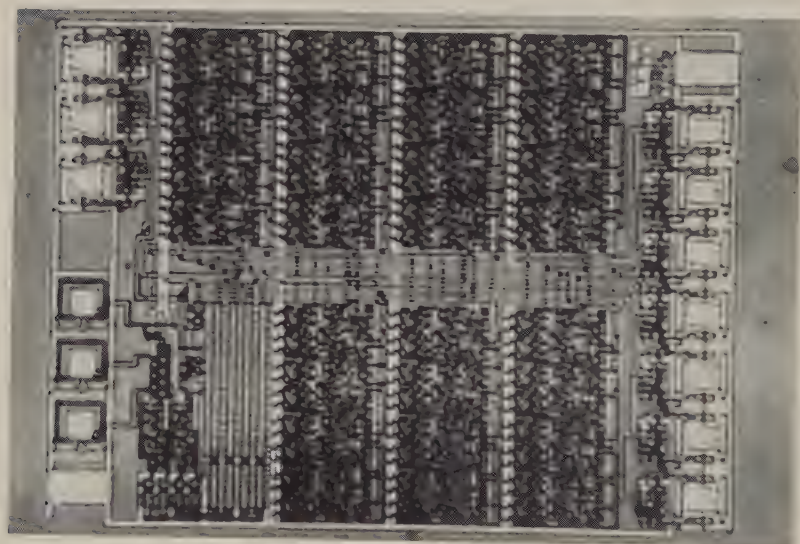


Fig. 88. Photograph of frequency divider.

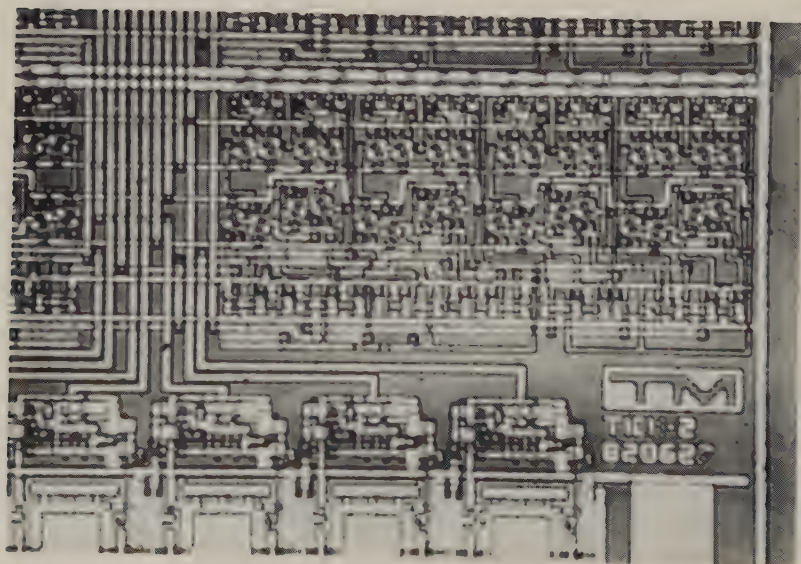


Fig. 89. Photograph of a decade counter part.

A small number of divider chips were mounted in ceramic packages to enable testing. A preliminary test was performed with the setup in fig. 90.

A sine wave signal from a function generator was fed to the input via a coupling capacitor and a bias network. The control input voltages were manually controlled by switches. Under normal operating conditions, i.e. 5 V supply voltage, 5 MHz input frequency and both control inputs low, the circuit worked. There were signals of correct frequency and duty cycle on all decade counter outputs. However, the pulse on the 1 PPS output was much shorter, approximately 10 ns, than the expected 100 ns. This was somewhat disappointing since the one pulse per second output is the most important. The short pulse duration is probably caused by the limited speed of the buffer stages in the standard output pad as suspected earlier (see appendix).

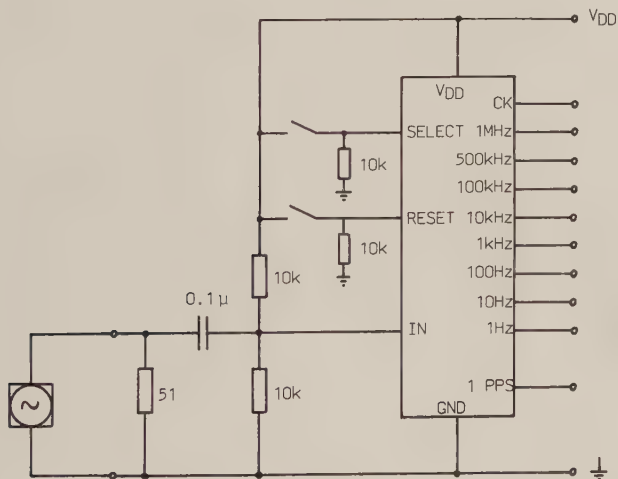


Fig. 90. Test setup.

The propagation delay times for a CMOS inverter depend on the transistor drain current I_d and the load capacitance. For a fixed load, the transition times can be reduced if I_d is increased. Fortunately there is a possibility to do this in spite of the fact that the chip already is processed, by increasing the supply voltage. At a V_{dd} of 7-8 V the output pulse length was in the vicinity of the theoretical value, and the overall performance was improved.

In spite of the higher supply voltage I was not able to make the divider work in the 10 MHz input mode. The first divider stage worked properly up to 20 MHz but the 1 PPS signal disappeared at frequencies higher than 7.5 MHz. This malfunction is probably caused by bad synchronization due to the simplified clocking scheme.

A number of test measurements were carried out to document the performance of the circuit. The results presented below are, in most cases, mean values of measurements on five separate chips.

The supply current was measured with a 1 V_{p-p} sine wave input signal and no external load on any output. The current is directly proportional to both supply voltage and input frequency.

f/V _{dd} :	5 V	7 V	10 V	:
5 MHz :	0.77	1.17	1.86	: mA
10 MHz :	1.46	2.26	3.54	: mA

The 1 PPS output pulse length was measured as a function of the supply voltage.

V _{dd} :	5	6	7	8	: V
t _p :	10	47	68	85	: ns

The results above suggest that 7 V would be a more suitable supply voltage than the originally intended 5 V.

The input sensitivity was estimated by reducing the amplitude of the sine wave input signal until the divider ceased to function.

Input sensitivity (V_{dd} = 7 V) : 50 mV_{p-p}

The transition times on the 1 PPS output were measured with the same input signal and a load of 10 kohm in parallel with 15 pF. As expected they decreased slightly with increasing supply voltage.

Transition times on 1 PPS output (V _{dd} = 7 V):	
Rise time :	20 ns
Fall time :	9 ns

The output drive current was measured by successively loading the output with resistors until the voltage had dropped 1 V.

Output drive current ($V_{dd} = 7 \text{ V}$):

```
-----
Ioh ( Voh = 6 V ) :    -2.4  mA
Iol ( Vol = 1 V ) :     6.0  mA
-----
```

Since all the outputs are identical, these values are assumed to be valid for any output.

The following delay times of interest were observed.

Total propagation delay times ($V_{dd} = 7 \text{ V}$):

```
-----
5 MHz input -> clock output :    44  ns
5 MHz input -> 1 PPS output  :    61  ns
-----
```

The jitter of the 1 PPS output pulses was investigated by time interval measurements between output pulses from two different test circuits with the same input signal. The standard deviation of 100 measurements was automatically calculated by the time interval counter. If the two circuits are assumed to contribute equally to the total jitter, the pulse jitter was found to be less than 0.5 ns.

A long term test was performed with two of the circuits on a prototype board. The outputs were compared against a tick signal from our cesium standard. During a test period of 40 days, not a single period of the input frequency was lost in any of the dividers. This implies that a very high reliability can be achieved under better conditions (better isolation, shielding etc.).

In conclusion, the frequency divider was not a complete success in all details. However, the performance with a 5 MHz input frequency is quite satisfactory and most design goals were achieved.

4.4 Tick generator module

A complete tick generator module (fig. 91) was designed to be used as a part of the passive hydrogen maser and for other time comparison applications. The complete circuit is contained in a small size shielding box.

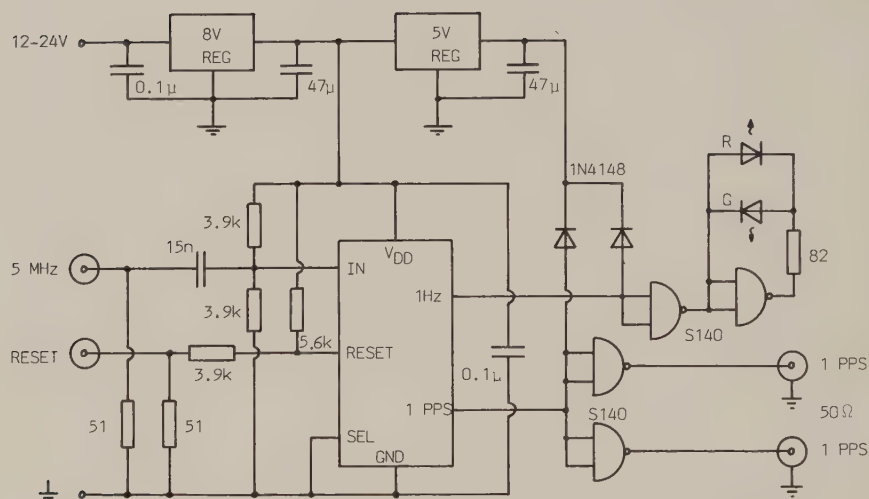


Fig. 91. Circuit diagram of tick generator module.

The 5 MHz and reset input are terminated in 50 ohms and the dc levels are adapted to the tick generator IC. The two separate tick outputs are buffered by 50 ohm line drivers. The symmetric 1 Hz square wave output is used to control two light emitting diodes indicating that the module is operating. The module also contains voltage regulators for the required supply voltages. The module accepts an input supply voltage of 12-24 V.

5 EXPERIMENTAL MEASUREMENTS

5.1 Magnetic field measurements

The most accurate method to measure the magnetic field in the storage bulb region is to use the magnetic field dependence of the hyperfine sublevels (see fig. 1). The energy difference between the two state selected levels ($F=1, m_F=0$) and ($F=1, m_F=+1$) may be calculated using the Breit-Rabi formula. For weak magnetic fields ($B < 10^{-5}$) the Zeeman transition frequency, f_z , is proportional to the magnetic field $\langle 14 \rangle$.

$$f_z = 1.400 \cdot 10^{10} \cdot B \quad (45)$$

where B = the magnetic field in Tesla

The Zeeman frequency is in the range from 1.4 kHz to 14 kHz for magnetic fields varying between 10^{-7} and 10^{-6} T.

To be able to measure the magnetic field in the storage bulb region a 16 turn Zeeman coil is placed outside the solenoid in the vicinity of the bulb (see fig. 92).

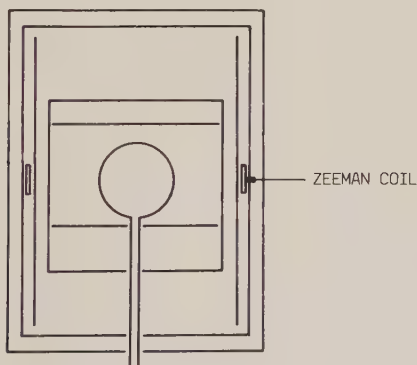


Fig. 92. Zeeman coil for magnetic field measurements.

The Zeeman coil is energized by a slow frequency sweep from a sine wave generator. When the generator frequency coincides with the Zeeman transition a dip in the 0.5 Hz

output amplitude will occur. To be able to operate at a higher signal level the hydrogen line servo loop is opened and the maser slightly detuned to obtain a larger 0.5 Hz modulation of the output signal. The sweep generator is controlled by a personal computer via the HP-IB interface bus and the output amplitude is recorded via an analog-to-digital converter and stored in a data file. The stored data from a sweep may be automatically analyzed and the magnetic field calculated from the center frequency of Zeeman resonance. The variations of the magnetic field within the bulb volume may be estimated from the linewidth of the recorded resonance. The results from one of the magnetic field measurements are shown in fig. 93.

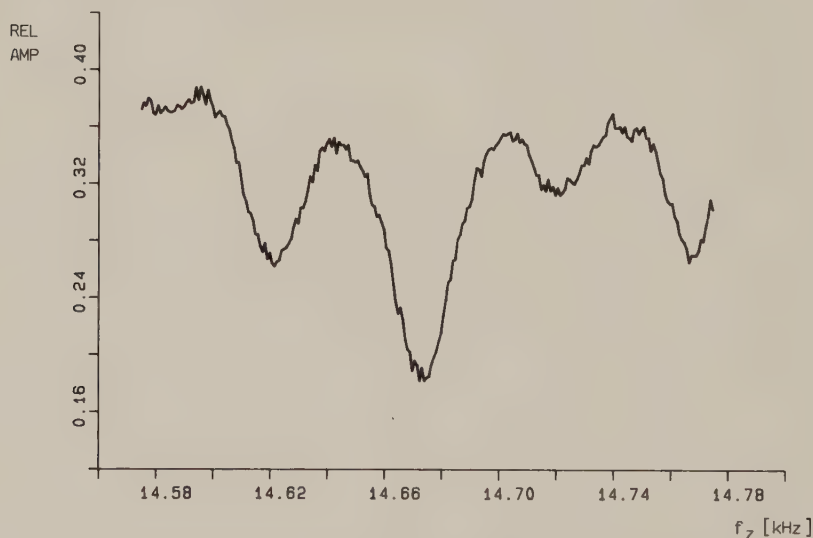


Fig. 93. Zeeman resonance from magnetic field measurement.

The output amplitude shows a distinct dip with a center frequency of 14673 Hz. The 'sidebands' 50 and 100 Hz from the main resonance are hum components from the mains.

The 14650 Hz Zeeman frequency corresponds to a magnetic field of $1.0481 \cdot 10^{-6}$ T. The measurements were repeated with the direction of the current in the solenoid changed, i.e. with the applied magnetic field reversed. These measurements gave a somewhat different resonance frequency, 15614 Hz corresponding to a magnetic field of $1.1153 \cdot 10^{-6}$ T. The difference is caused by residuals of the external magnetic fields, mainly the earth's magnetic field.

The magnitudes of the applied and residual magnetic fields may be calculated from the results of the two measurements:

Applied magnetic field : $1.0817 \cdot 10^{-6}$ T
Residual magnetic field: $3.36 \cdot 10^{-8}$ T

The linewidth of the resonance is approximately 25 Hz, corresponding to variation in the magnetic field of $1.8 \cdot 10^{-9}$ T. This indicates that the uniformity of the magnetic field in the storage bulb region is better than 0.2%.

The solenoid current used during the measurements was 0.5009 mA. Since the magnetic field is proportional to the coil current, a coil factor can be calculated as the ratio between the magnetic field and the current. This factor may be used to adjust the magnetic field to a desired value by setting the coil current.

$$B = 2.1595 \cdot 10^{-6} \cdot I \quad (46)$$

where B = the magnetic field in Tesla
 I = the solenoid current in mA

5.2 Hydrogen resonance linewidth measurements

The linewidth of the hydrogen resonance can be estimated by studying the phase of the 0.5 Hz amplitude modulation while the microwave frequency is swept slowly across the hydrogen line. For this purpose the hydrogen line servo loop is opened and the servo amplifiers changed from integrators to unity gain amplifiers by connecting resistors in parallel with the feedback capacitors. Since the quartz oscillator is not locked to the hydrogen transition in this operating mode, it is replaced by a more accurate standard (cesium or rubidium) during these measurements. The slow frequency sweep is generated by varying the 19.6 MHz synthesized signal.

Since the required sweep time is very long (2-10 hours), due to the high Q of the hydrogen resonance and the time constant of the servo amplifiers, the built-in sweep facilities of the synthesizer cannot be used. Therefore the synthesizer output frequency is controlled via the HP-IB interface by a personal computer. A sweep control program was written where the center frequency, the sweep width and the sweep time may be selected independently. The synthesizer frequency is adjusted once every second by an amount determined from the sweep width and the sweep time.

At intervals depending on the sweep time the computer reads the output amplitude from the synchronous detector via an analog-to-digital converter input and stores the results in a data file. Sweep data and a graph of the recorded response are presented on a colour video display during the sweep so that the progress of the measurement can be monitored. Fig. 94 shows the typical S-shape of the phase detector response.

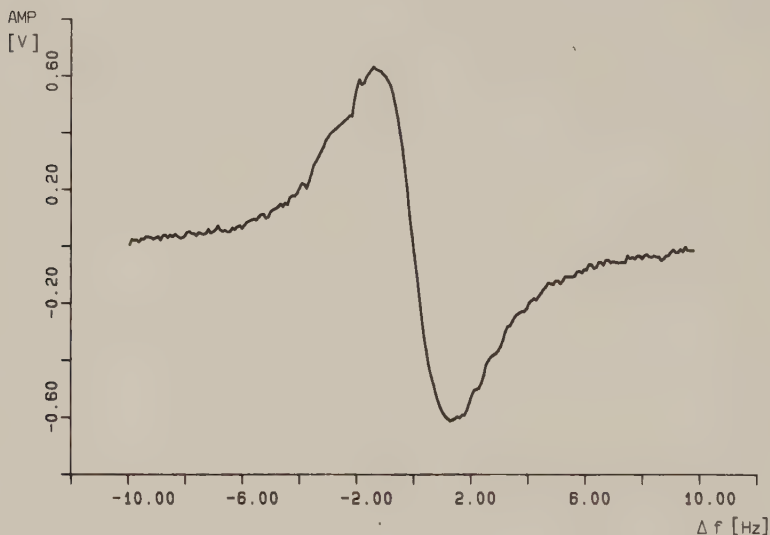


Fig. 94. Typical phase detector response.

After the completion of a sweep measurement a second program may be used that plots the response in more detail and calculates the maximum and minimum amplitudes. A preliminary estimation of the linewidth is given by the frequency difference between the maximum and minimum amplitude points. Finally, the slope of the curve at the zero crossing is calculated.

The phase detector response depends not only on the shape and linewidth of the hydrogen resonance but also on the frequency deviation caused by the 0.5 Hz phase modulation of the microwave signal. To be able to draw more exact conclusions regarding the linewidth from the results of the sweep measurements, a computer program was written that calculates and plots a theoretical phase response for a specified linewidth and frequency deviation.

The amplitude response from a Lorenz-shaped resonance is calculated for one period of the modulating signal. The resulting amplitude depends on the offset of the center frequency of the interrogating frequency sweep compared to the center frequency of the hydrogen resonance. The phase of the amplitude response is converted to an output voltage in a mathematical model of the synchronous detector. The complete phase detector response is generated by moving the center frequency of the interrogating sweep across the Lorenz resonance curve. Fig. 95 to 99 show the theoretical phase detector responses for various combinations of resonance linewidth B and frequency deviation D .

The width of the response (the difference between the frequencies at the minimum and maximum amplitudes) decreases as the frequency deviation is reduced. The amplitude of the response has a maximum at a deviation slightly greater than half the linewidth. The slope of the response, which is the most important factor when the crystal oscillator is locked to the hydrogen line, has its maximum when the deviation is approximately half the linewidth.

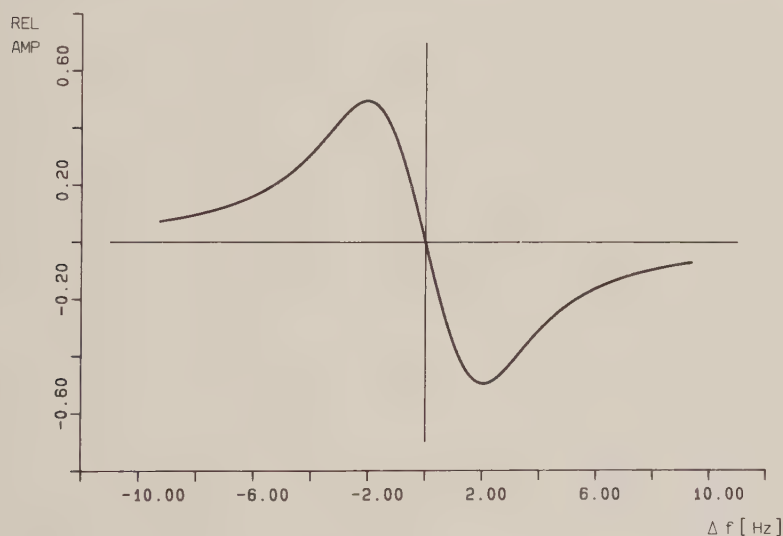


Fig. 95. Phase detector response - $B=4$ Hz and $D=2$ Hz

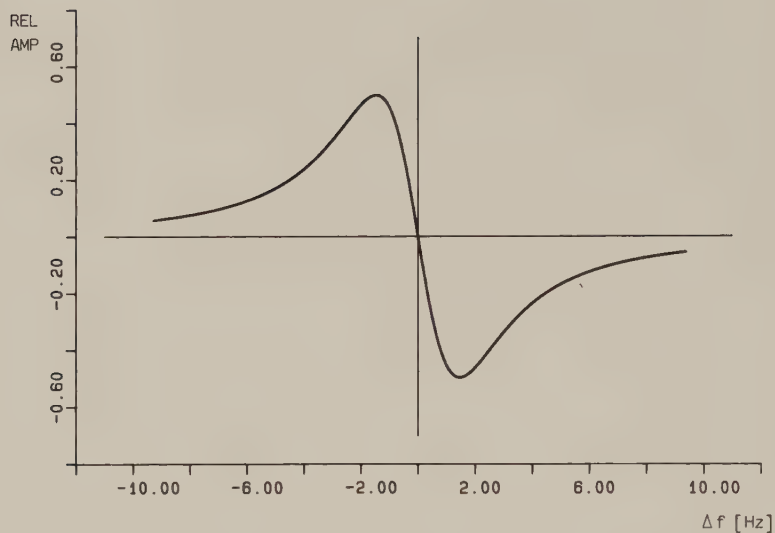


Fig. 96. Phase detector response - $B=4$ Hz and $D=0.5$ Hz.

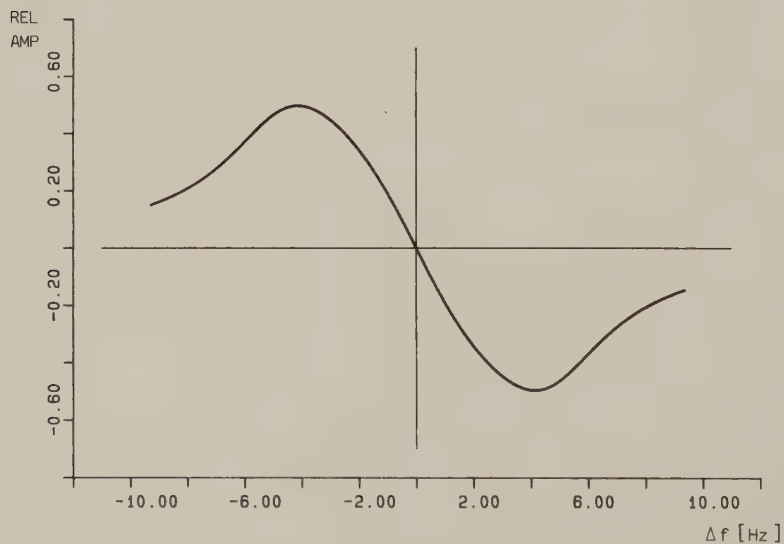


Fig. 97. Phase detector response - $B=4$ Hz and $D=5$ Hz.

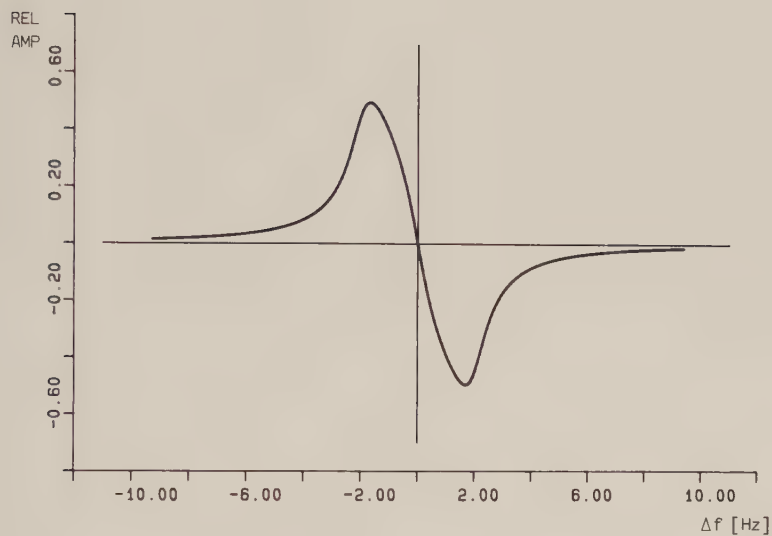


Fig. 98. Phase detector response - $B=1$ Hz and $D=2$ Hz

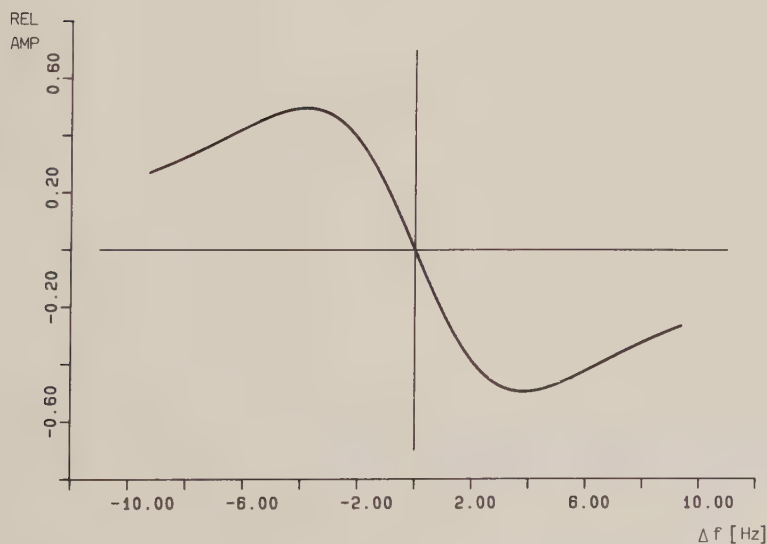


Fig. 99. Phase detector response - $B=10$ Hz and $D=2$ Hz

The conclusion is that the frequency deviation should equal half the linewidth to achieve optimum locking conditions. In that case the frequency difference between the maximum and minimum points is a fairly good estimation of the half-power linewidth of the hydrogen resonance. To further increase the accuracy of the estimation the computer program was supplemented with an optimizing facility that fits a theoretical response to the measured response. From the frequency deviation used at the measurement the program estimates the actual linewidth and plots both responses in a diagram. Fig. 100 and 101 show the results of two linewidth measurements using a 0.25 Hz frequency deviation. The sweep width was 20 Hz with a sweep time of 10 hours in both measurements.

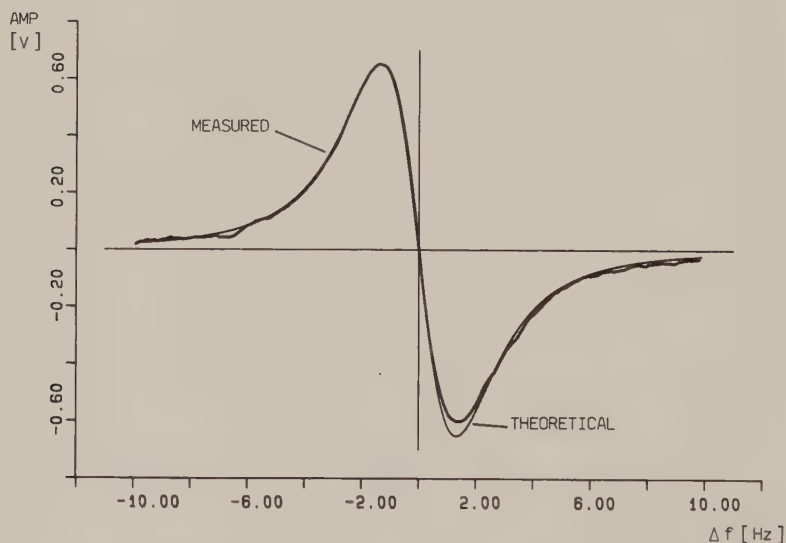


Fig. 100. Result of linewidth measurement - 1.

The fit is very good in both cases and both resulted in an estimated linewidth of 4.5 Hz. The total amplitude of the response is 1.25 V and the slope at the zero crossing is -0.9 V/Hz (minimum 0.5 Hz gain).

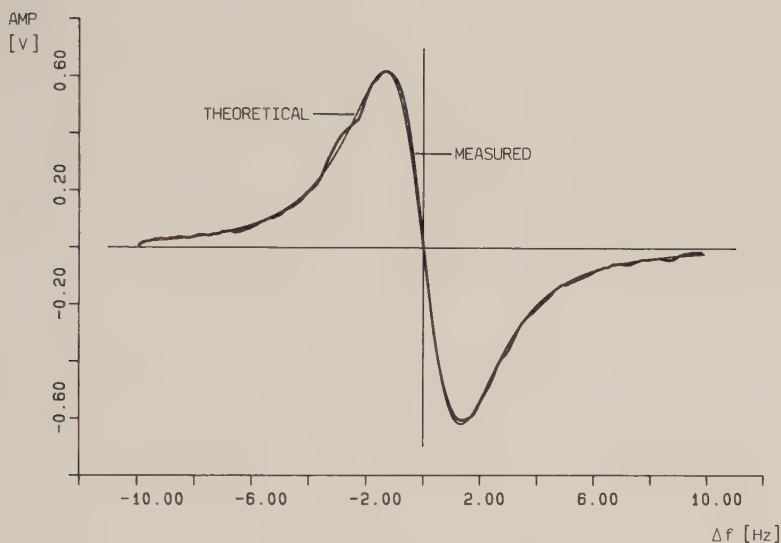


Fig. 101. Result of linewidth measurement - 2.

5.3 Frequency stability measurements

The long-term stability of the passive hydrogen maser frequency standard was compared to a cesium standard using time interval measurements (fig. 102).

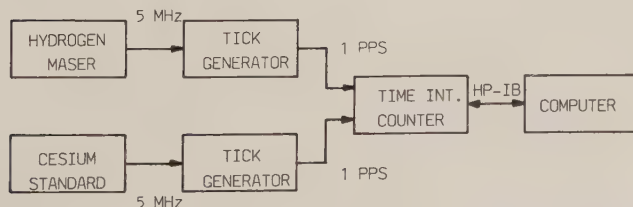


Fig. 102. Long-term stability measurement.

The 1 PPS signals are derived from the 5 MHz output frequencies of the hydrogen maser and the cesium standard using two of the designed tick generator modules. The time intervals between the tick pulses are measured by a time interval counter. The counter is controlled by a personal computer that initiates the measurements at regular intervals and stores the results for future evaluation.

The stability is calculated from measurement results using the Allan variance $\sigma_y^2(\tau) < 1 >$.

$$\sigma_y^2(\tau) = \frac{1}{N-1} \sum_{k=1}^{N-1} \frac{(\bar{y}_{k+1} - \bar{y}_k)^2}{2} \quad (47)$$

where τ = the averaging time for each sample

N = the number of sample averages

$\bar{y}_k$ = sample averages

Each sample average $\bar{y}_k$ of the frequency stability is derived from two consecutive time interval measurements Δt_{i+1} and Δt_i .

$$\bar{y}_k = \frac{\Delta t_{i+1} - \Delta t_i}{T} \quad (48)$$

where T = the repetition time = the sample time

The following long-term stability is the result of measurements during one week with the maser locked to 1420 405 751.934 Hz.

$$\sigma_y(\tau) = 3.9 \cdot 10^{-11} \quad \tau = 1 \text{ hour} \quad (49)$$

$$\sigma_y(\tau) = 2.8 \cdot 10^{-12} \quad \tau = 1 \text{ day} \quad (50)$$

In addition the phase difference between a 1 MHz signal from the hydrogen maser and a 1 MHz signal from the cesium standard was monitored on a strip chart recorder to obtain further information about the stability.

6 CONCLUSIONS AND FURTHER RESEARCH

6.1 Estimation of total accuracy

In order to estimate the total accuracy of the designed passive hydrogen maser, the frequency shifts from different perturbations must be calculated.

- The frequency shift due to cavity pulling is eliminated by the cavity servo.
- The frequency shift Δf_M caused by the applied magnetic field of $1.0481 \cdot 10^{-6}$ T is given by equation (20):

$$\Delta f_M = 2750 \cdot 10^8 \cdot (1.0481 \cdot 10^{-6})^2 = 0.302 \text{ Hz}$$

- The frequency shift due to the Crampton effect is on the order of 10^{-13} and is neglected here.
- The second order Doppler shift Δf_D may be calculated from equation (21) for a bulb temperature of 41.5°C (314.5 K):

$$\begin{aligned} \Delta f_D &= - \frac{3 \cdot 1.381 \cdot 10^{-23} \cdot 314.5}{2 \cdot 1.673 \cdot 10^{-27} \cdot (2.998 \cdot 10^8)^2} \cdot 1.420 \cdot 10^9 = \\ &= - 0.062 \text{ Hz} \end{aligned}$$

- The frequency shift Δf_W caused by wall collisions is, for a 14 cm bulb, approximately (eq. (22)):

$$\Delta f_W \approx - \frac{355 \cdot 10^{-3}}{14} = - 0.025 \text{ Hz}$$

- The frequency shift due to spin exchange is a few parts in 10^{-13} for a passive maser and may be neglected here.

The total calculated frequency shift Δf is

$$\Delta f = 0.302 - 0.062 - 0.025 = 0.215 \text{ Hz}$$

which results in a theoretical operating frequency of

$$f = f_0 + \Delta f = 1402\,405\,751.984 \text{ Hz}$$

The operating frequency that corresponds to minimum drift compared to cesium is measured as

$$f = 1402\,405\,751.960 \text{ Hz}$$

The 0.024 Hz difference is equal to a fractional frequency accuracy of $1.7 \cdot 10^{-11}$.

6.2 Areas for future improvements

In the effort to achieve increased accuracy and long-term stability of a frequency standard there is almost no area in which improvements cannot be made. The passive hydrogen maser frequency standard described in this work has only been operating for a short period of time and many problems remain unsolved. Some of the most obvious improvements in different areas are listed below.

Hydrogen source:

- Control of the hydrogen pressure in the dissociator and the rf discharge power in order to achieve more constant beam intensity, thus reducing the variations of the spin exchange shift which depend on hydrogen density.
- Optimization of the hydrogen beam intensity to reduce the linewidth of the hydrogen resonance.

State selector:

- Adjustment of the beam geometry in order to improve the efficiency of state selection.
- Installation of a beam-stop of suitable size to block the unselected atoms on the center axis of the hexapole magnet.

Vacuum system:

- The cavity should be evacuated by the roughing system to reduce the influence of the environment on the cavity resonance frequency, thus reducing the demands on the cavity servo.

Storage bulb:

- Use of a larger storage bulb would reduce the volume dependent wall shift.
- Experimentation with other shapes of storage bulb, e.g. cylindrical.
- Improvement of the temperature control of the storage bulb through better insulation and perhaps installation of an outer oven to reduce the influence of room temperature.

Microwave cavity:

- Filling the space between the outer and inner end plates of the cavity with a microwave absorbing material to reduce the influence of changes in the electromagnetic environment on the resonance.
- Experimentation with other loop sizes to optimize the coupling to the hydrogen line.
- Experimentation with other cavity materials.

Magnetic field:

- Improvement of shielding to allow operation of the maser at a weaker magnetic field.
- Measurement of the variations of the magnetic field in the storage region to evaluate the performance of the constant current source.
- Improvement of the magnetic field measurement technique to eliminate the hum components.

Microwave signal generation:

- Investigation of the amplitude caused by the tuned multiplier chain and its influence on the cavity (and hydrogen line) servo.
- Reduction of the lengths of the cables carrying microwave signals by moving the rf section closer to the microwave cavity.
- Improvement of the shielding of rf cables to reduce crosstalk.
- Insertion of isolators at the input and output loops of the cavity to avoid detuning due to changes in the reactive load of the cavity.
- Measurement of the variations of the microwave power level.

Cavity and hydrogen line servos:

- Investigation of the cause of the fluctuations of the cavity fine tuning voltage which exist in spite of temperature control of the cavity.
- Optimization of the time constants of the cavity and hydrogen line servos.
- Optimization of the microwave cavity input signal level to obtain the best locking condition.

Measurements:

- Measurement of the hydrogen atom lifetime by studying the transient behavior of the maser.
- Further measurements of the short-term and long-term stability of the passive maser.

APPENDIX

1 DETAILED DESCRIPTION OF TICK GENERATOR IC DESIGN

1.1 Flip-flop design

The master-slave toggle flip-flop consists of inverters and transmission gates, which are the two most commonly used basic building blocks in CMOS digital integrated circuits. The schematic diagram symbols for these circuits are shown in fig. 103.

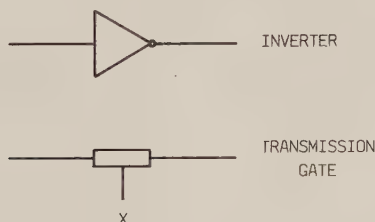


Fig. 103. Symbols for inverter and transmission gate.

Both the inverter and the transmission gate contain only two transistors, one n-channel and one p-channel. Since the gate of the MOS transistor is electrically isolated from the channel between source and drain by the gate oxide, the input impedance of an inverter is purely capacitive. This property makes it possible to use the inverter as a dynamic memory cell. If the input capacitor is charged by a logical signal level which is later disconnected, the input voltage will remain unchanged for a period of time. This storage time is dependent on the size of the input capacitor and the magnitudes of leak currents. A transmission gate is used to connect and disconnect the input signal to the inverter. When the control signal X is low both transistors are cut off and the impedance between input and output is very high. When X is high the transistors are on and there is a low

impedance connection between the input/output terminals. N-channel and p-channel transistors are used in parallel to ensure a low on-resistance independent of the signal level. Under most circumstances the transmission gate can be considered to be an almost ideal bidirectional switch.

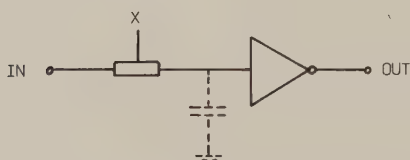


Fig. 104. Dynamic memory cell.

The dynamic memory cell (fig. 104) consists of only four transistors and therefore it has small dimensions. The storage time of a dynamic memory cell is, however, limited to some milliseconds. In systems with clock frequencies in the MHz-range this is quite sufficient. If you want to store the signal for a longer period, independent of the electrical time constants in your system, the cell can be made static by introducing positive feedback (see fig. 105).

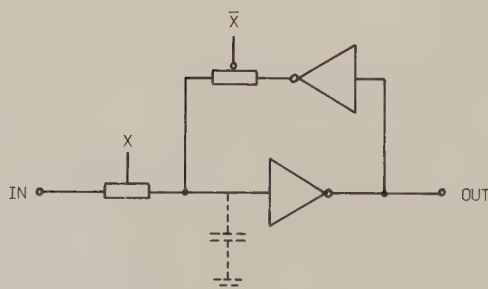


Fig. 105. Static memory cell.

The feedback path contains an additional inverter and

transmission gate. When the load signal X is high the input signal is connected to the inverter input and the feedback loop is opened by the second transmission gate. As the load signal goes low, the input signal is disconnected and feedback loop closed. The input level of the first inverter is held constant by the second inverter. Thus the output signal is constant until another value is loaded. A disadvantage of the static memory cell is that the number of transistors needed is doubled.

The type of memory cell described above is often referred to as a latch. It can be noticed that when the load signal is high, the input signal is fed directly to the output. The latch is said to be transparent. This property is a problem if you want to use the cell in a shift register or in a counter, but it can be solved by using a master-slave configuration.

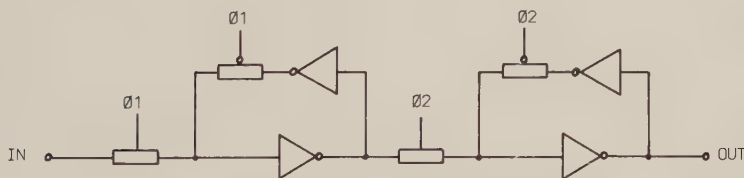


Fig. 106. Static master-slave latch.

The static master-slave latch (fig. 106) consists of two cascaded static memory cells. During the first phase of the clock signal the input signal is fed to the first latch (the master) and the input of the second latch (the slave) is disconnected from the output of the first. Thus a change on the input does not affect the output at this stage. During the second phase the master is isolated from the input and the content of the latch is transferred to the slave and thus to the output. It is important that the two phases do not overlap, since that would cause the output to be directly connected to the input for a short time.

The latch can be changed into a toggle flip-flop by introducing feedback from the outputs to the the input. If the inverted output $\bar{Q}$ is connected to the input via a transmission gate the flip-flop will toggle each time a clock pulse is encountered. If on the other hand the feedback is taken from the noninverted output Q , the content will not change when the flip-flop is clocked. A synchronous

reset is obtained if the feedback path is broken and the input is connected to a logical zero independent of the previous status. A complete synchronous, static master-slave flip-flop with synchronous reset can be of the design in fig. 107.

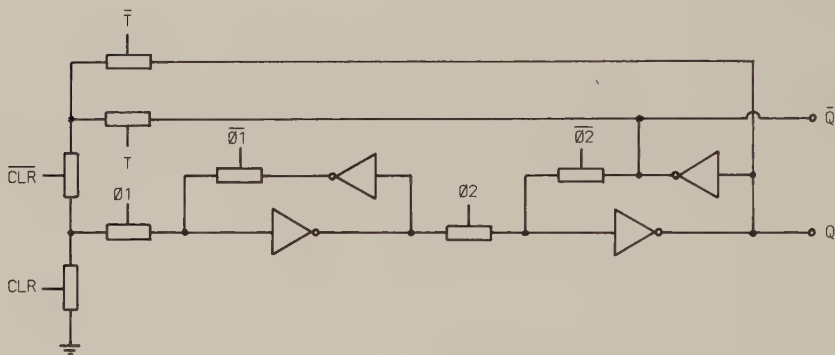


Fig. 107. Synchronous static master-slave flip-flop with reset.

To avoid direct connection between input and output, the clock pulses ϕ_1 and ϕ_2 must not be overlapping. The transmission gates must also have access to the inverse of the control signal. This means that, besides ϕ_1 and ϕ_2 , $\bar{\phi}_1$ and $\bar{\phi}_2$ must also be distributed to all flip-flops (if you do not want to add an inverter to each transmission gate). If it were possible to use $\bar{\phi}_1 = \phi_2$ and $\bar{\phi}_2 = \phi_1$, and thus reduce the number of clock signals, the design of the flip-flops could be made far less complex. Since this also results in a smaller clock distribution network, it was decided to try this concept to reduce the chip area requirement.

1.2 Simulation and breadboarding of basic circuits.

Once the chip is processed, it is impossible to make any changes in the construction. If the circuit does not work properly, the only way to correct it is to make the necessary changes in the original layout and reprocess it. This is a very costly procedure so it is preferable to minimize repetitions. Therefore you should try to verify on

beforehand that the design will work as it was intended.

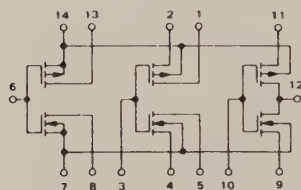
There are basically two different methods of circuit verification:

- 1) Construction of a prototype from standard components on a breadboard.
- 2) Circuit simulation by a computer program.

In the breadboard method you build up your construction with standard components or, if it is possible, with transistors and circuit elements manufactured in the process you are going to use. If the circuit is complex and consists of a large number of transistors, it is not practicable to breadboard the complete circuit, only the most critical parts. Another great disadvantage is that you cannot get a realistic test of fast courses of events due to the comparatively large stray capacitances that are introduced by the connecting wires.

The second and more often used method is computer simulation. Transistors and other components are described by mathematical models which are used to calculate the behavior of the circuit. Macro models can be used to describe larger function blocks like amplifiers etc. The accuracy of the simulation result is highly dependent on correct input values of the parameters in the transistor models. A model for an MOS-transistor may contain over 25 different parameters, all dependent on the process used. Correct estimation of parameter values from static or dynamic measurements on test structures is essential to be able to perform useful simulations. The execution time for the simulation program puts a practical limit on the number of components in the simulated circuit. At that time we only had access to SPICE, which is the most commonly used numerical simulation program.

Since the design of flip-flops using inverters and transmission gates was a new concept, it was necessary to test the basic toggle flip-flop on a prototype board. The necessary MOS-transistors were taken from a standard CMOS circuit MC14007UB (fig. 108) which contains two complementary transistor pairs and an inverter. This circuit is manufactured in a different process but it can be used to verify the logical function of the flip-flop design. The flip-flop shown in fig. 97 was built (except for the reset gates) and operated with a supply voltage of 5 volts and with an external clock signal from a function generator. Measurements showed that it worked properly with a propagation delay time from clock input to the noninverting output of about 240 ns. Maximum operating frequency was approximately 2 MHz - as high as could be expected with the transistors used and the additional capacitances outside the chip.



VDD = Pin 14
VSS = Pin 7

Fig. 108. Schematic diagram of MC14007UB.

To investigate the behavior at higher frequencies and with SOS-transistors, the same circuit (with reset) was simulated with SPICE. Parameter values for the SOS/CMOS process were supplied by the manufacturer. Several simulations were performed with varying clock frequencies and clock phase relations. All n-channel and p-channel transistors in inverters and transmission gates were of minimum size, i.e. length = 4 μm and width = 5 μm . One of the simulations was made with an input frequency of 20 MHz and clock pulses with rise- and fall-times of 2.5 ns. Only two clock signals, $\phi 1$ and $\phi 2$ in fig. 109, were used as suggested earlier. The simulated waveforms are shown in fig. 110.

The following approximate transition times were obtained:

Inverter input master:	rise-time	14 ns
	fall-time	6 ns
Inverter input slave :	rise-time	6 ns
	fall-time	3 ns
Non-inverting output :	rise-time	3 ns
	fall-time	2 ns

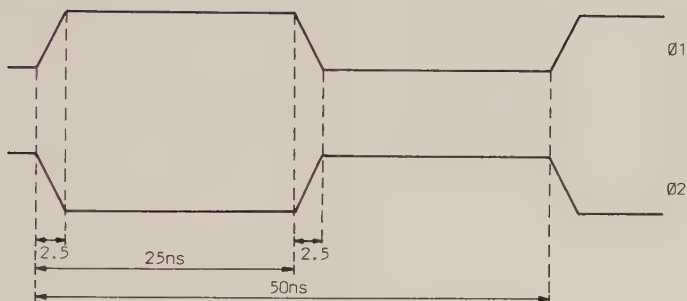


Fig. 109. Clock signals in SPICE simulation.

The results from this simulation indicated that the flip-flop design with the simplified clock scheme should work as expected and at frequencies well above the needed 10 MHz. It should, however, be noticed that these results are only valid for on-chip conditions with small load capacitances. Special care must be taken to maintain high operating speed for signals driving many gates and output signals. Fig. 110 shows the simulated signal waveforms at some interesting nodes in the flip-flop circuit.

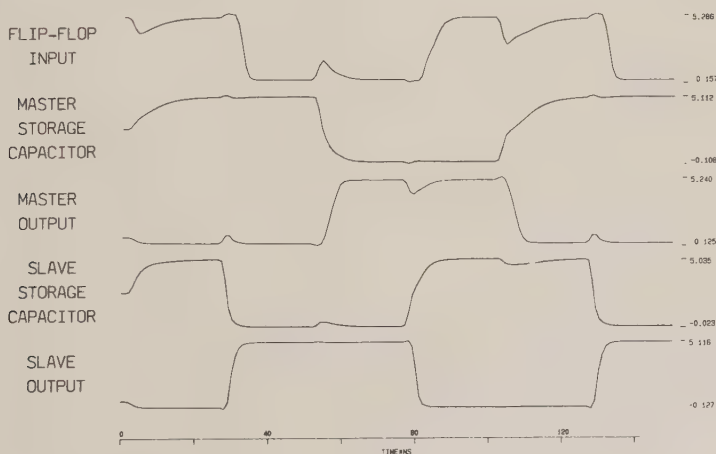


Fig. 110. Simulated waveforms.

A complete decade counter was also breadboarded to test the counting sequence and carry generation. J-K flip-flops and logic gates from the MC14000 CMOS family were used as components. Finally, the "discrete" counter was cascaded with integrated synchronous decade counters MC14162B to verify that it worked properly connected to other counters in a longer chain.

After these checks it was more probable that the logic designs of the flip-flop and the counter were correct and that it was possible to achieve the desired counting speed with the type of clock signals chosen.

1.3 Layout generation

The layout aid used to design the circuit is an Apple II personal computer with an LSI design software package developed at the Department of Computer Engineering, Lund Institute of Technology <45>. This system is based on the Mead-Conway method of structured VLSI design <43>. The geometric patterns of the different mask layers are described by a program written in Pascal, supplemented with a set of basic layout functions like drawing wires and boxes, defining the current layer, etc. This allows you to generate dynamic geometries that are dependent on various parameters.

By means of structured programming the description of the layout can be made relatively compact, especially if the circuit design itself is well structured. A set of fundamental cells like inverters, transmission gates etc. are combined into more complex cells. In this way the complete chip is built up with cells on different levels. It is a great advantage if most of the connections between the functional blocks are contained within the cells themselves. This may substantially reduce the routing problems.

The floor-plan in fig. 111 was found to give a regular chip layout with short signal paths, particularly from the 5 MHz input to the 1 PPS output. It is also important that the interconnections between different blocks and the distribution of supply voltages be considered at an early stage to avoid complications later.

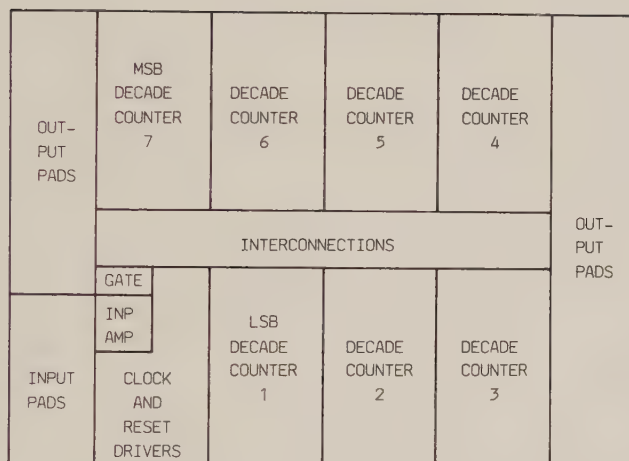


Fig. 111. Floor-plan of tick generator circuit.

The decade counters (fig. 112) consist of four identical flip-flops and logic circuitry to generate the toggle conditions that provide the proper counting sequence.

Clock and reset pulses are common to all flip-flops. Therefore there must be vertical bus lines through the cells for these signals and for the supply voltage. All inputs and outputs are connected to the logic circuitry and must be placed on the right side of the cell. Connections between the counters are situated in the interconnection area in the center of the chip, which implies that the inputs and outputs of the complete decade counter must be placed on the upper side of the counter cell. For the same reason the counters in the upper half of the chip are mirrored across the horizontal axis.

The three input pads for the input frequency, reset signal and select signal are situated in the lower lefthand corner of the chip. The input amplifier is placed close to the 5 MHz input. Both clock and reset pulses are converted into two-phase signals and connected to the counters via buffer circuits.

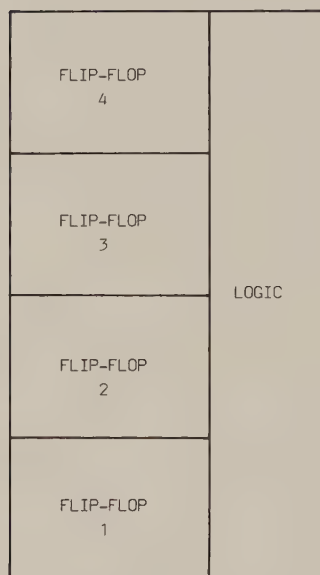


Fig. 112. Structure of decade counter cell.

The decade counter chain is arranged in a folded manner so that the gate connecting the clock signal to the output during one period each second can be put close to both the most significant decade and the input amplifier. The gate output is then connected to the nearest output pad. This configuration is used to minimize the signal path from the 5 MHz sine wave input to the 1 PPS output, which will result in lower pulse jitter. The remaining square wave outputs from the different decade counters are connected to the nearest available output pad without any special considerations.

After this general planning of the total chip area and positioning of the different function blocks it is time to proceed to the more detailed layout of the basic cells. The cell that occurs most frequently on the chip is the flip-flop. It is repeated 28 times, which makes it worthwhile to put some effort into making it as compact as possible.

The main object of the layout work is to supply the manufacturer with basic data for mask generation. The exact geometry of all mask layers must be specified in detail. In the layout program each logical layer is represented by a specific colour on the graphics display. The following logical layers are used:

- Green: N-doped silicon for n-channel transistors.
- Yellow: P-doped areas for p-channel transistors.
- Red: Poly-silicon for transistor gates and wires.
- Blue: Metal for wires.
- Cut: Contact holes in thick oxide.
- Glass: Holes in glass protection for bonding areas.

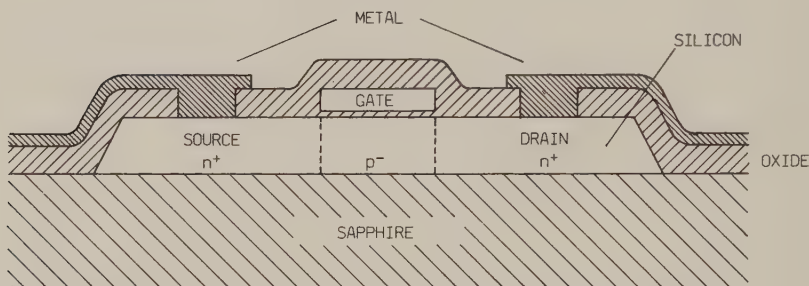
All layout dimensions are given in standard units for highest flexibility. For the SOS/CMOS-process used 1 unit corresponds to 1 μm at present time. All process limitations are simplified into a set of design rules <44> which guarantee that the chip works if they are kept. These rules specify, among other things, minimum widths of wires (red and green 4 units, blue 5 units), minimum distance between wires (red and green 4 units, blue 5 units), minimum contact size (5x5 units), minimum transistor dimensions (width 4 units, length 5 units), and various safety distances.

To be able to understand the layout of the simple cells, it might help to look at the structure of a single n-channel MOS-transistor and the corresponding graphic layout (fig. 113).

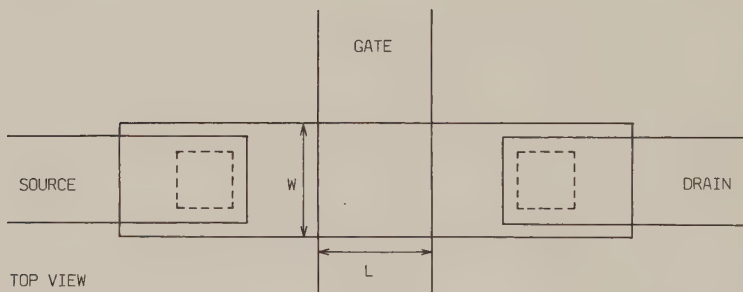
The transistor is made from a small "island" of silicon on the insulating sapphire substrate. The polysilicon gate is placed over a weakly p-doped channel area between the strongly n-doped source and drain areas. It is insulated from the other parts of the transistor by a thin layer of silicon oxide. In layout terms it means that an n-channel transistor is generated wherever green and red cross. If the green is surrounded by a yellow box, the transistor becomes a p-channel type. The blue wires are metal connection leads to source and drain.

The n-channel MOS-transistor works as follows. If the gate-source voltage (V_{gs}) is lower then the so-called threshold voltage ($V_t = 0.7 - 0.9 \text{ V}$), there is no electrical connection between source and drain. The transistor is said to be cut off. For higher gate-source voltages the region under the gate is inverted into an n-conducting channel. Now a drain current (I_d) can flow from

drain to source with a magnitude depending on how much the gate-source voltage exceeds the threshold voltage.



SIDE VIEW



TOP VIEW

Fig. 113. Structure of n-channel MOS-transistor.

The output characteristics of the MOS-transistor can be divided into two regions depending on the drain-source voltage (V_{ds}). Simplified expressions for the drain current are presented below.

Linear region ($V_{ds} < V_{gs} - V_t$):

$$I_d = C_{ox} \cdot \mu \cdot \frac{W}{L} [(V_{gs} - V_t) \cdot V_{ds} - \frac{V_{ds}^2}{2}] \quad (51)$$

Saturation region ($V_{ds} > V_{gs} - V_t$):

$$I_d = C_{ox} \cdot \mu \cdot \frac{W}{L} \cdot \frac{(V_{gs} - V_t)^2}{2} \quad (52)$$

where C_{ox} = oxide capacitance
 μ = mobility
 W = transistor width
 L = transistor length

In digital applications the transistors are either cut off or saturated. The drain current in the on-state is (for a certain supply voltage) then only depends on the physical dimensions (W/L). The same equations apply for p-channel transistors but the mobility is lower by a factor of two, which results in less drain current for a p-transistor of equal dimensions.

The CMOS inverter (fig. 114) is a pair of complementary transistors with gates and drains tied together.

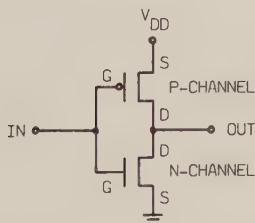


Fig. 114. CMOS inverter.

Since there is no connection to the substrate for SOS-transistors, the simplified transistor symbols above are used in all circuit diagrams.

When the input signal is high, i.e. equal to V_{dd} , the n-channel transistor is turned on and the p-channel transistor turned off, which makes the output to go low. If the input is high, the situation is reversed and the output goes high. Thus the output signal always is the inverse of the input signal. The transition time, i.e. the time needed to change from high to low level (or vice versa), depends on the load impedance. Normally the inverter is loaded with one or more identical inverters. The input impedance of the

inverter is the input impedances of the two transistors in parallel. Since the gate is insulated from the other parts of the transistor, the input impedance of the transistor consists only of capacitances. The two main contributions to the gate capacitance are the overlap capacitance to source and drain and the capacitance to the channel (if the transistor is on).

The transition times for an inverter with minimum size transistors loaded with one minimum inverter input capacitance can be estimated by simple calculations:

$$t(\text{high-low}) = 1 \text{ ns}$$

$$t(\text{low-high}) = 2 \text{ ns}$$

The transition time from low to high is longer due to the lower output conductance of the p-channel transistor. This fact must be kept in mind when maximum switching speed is required.

The layout of the inverter can be varied in many ways depending on the demands of the particular application. For use in the flip-flop the folded layout in fig. 115 was found to be the best in conjunction with the transmission gates.

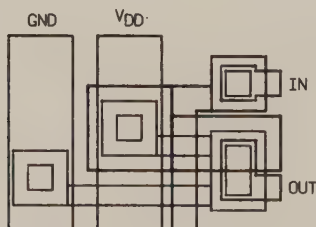


Fig. 115. Layout of inverter.

The connections to ground and V_{dd} are placed as vertical bus wires to the left in the cell. Both input and output are accessible on the right side of the boundbox. This allows several inverters to be stacked on top of each other with their supply voltage wires automatically connected. All dimensions are made as small as possible according to the design rules except for the supply leads which are made wider (13 units) to reduce the voltage drop at high currents.

As mentioned earlier the MOS-transistor may (in digital applications) be regarded as an almost ideal switch controlled by the gate voltage. This property is also very useful when you want to build a bidirectional analog switch to be put in series with the signal. To ensure that the on resistance is low for all signal levels, one n-channel and one p-channel transistor are connected in parallel. Such a circuit is called a transmission gate (fig. 116).

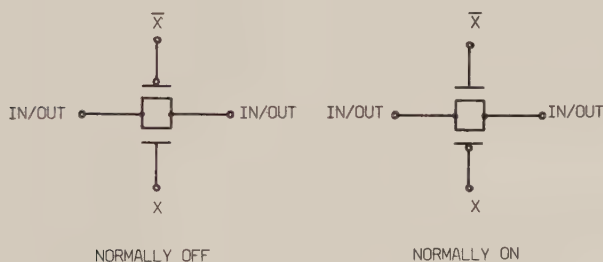


Fig. 116. Transmission gate.

The switch is controlled by the control signal X and its inverse $\bar{X}$. The switch is on if X is high and off if X is low. The switch is said to be normally off. It can be changed into a normally on switch by letting the transistors change place without altering the control signals. Since both types are frequently used, it would be advantageous to design a transmission gate cell that could be used as either one of them.

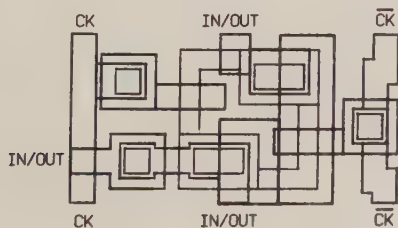


Fig. 117. Transmission gate layout.

The transmission gates in the flip-flop are controlled by the same clock and clear signals. Therefore the transmission gate cell must be made cascadable and have the control signal wires running through the cell. It is also desirable that this cell have approximately the same height as the inverter cell. These demands resulted in the layout in fig. 117.

The inputs/outputs are in the center of the top and bottom sides of the cell. This allows the the cell to be mirrored and/or rotated before it is cascaded. An additional input/output is placed on the left side to enable access to the signal between two cascaded transmission gates without having to separate them. By specifying a parameter in the procedure call in the layout program the yellow area is placed in two different positions. In this way you can select whether the switch is to be normally on or normally off.

The next step is to combine a number of the two previously designed cells in the flip-flop layout. As an aid in configuring the basic cells so as to minimize the routing between their respective inputs and outputs, you can use a simplified layout diagram called a stick diagram. These diagrams, where all wires are represented by single strokes, are very helpful when solving topology problems. Most of the wiring to connect the different cells is made with red and blue wires since they can be crossed without generating any transistors or causing electrical contact. To avoid too many conflicts it was tried to systematically use red for the horizontal wires and blue for the vertical throughout the whole design. After some puzzle work the result illustrated in fig. 118 was achieved.

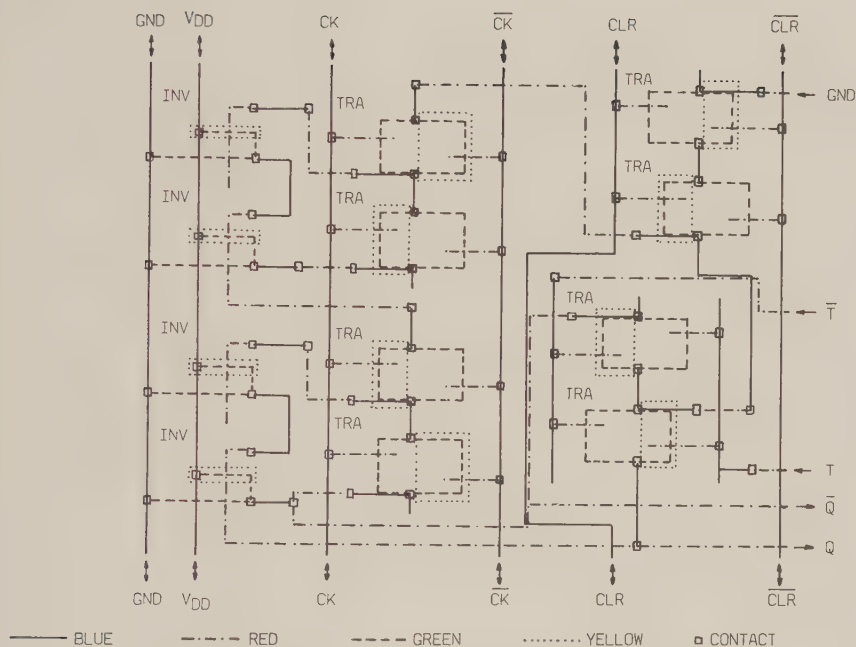


Fig. 118. Stick diagram of flip-flop cell.

As earlier planned, supply voltages and clock and clear signals with their inverses are distributed to the concerned cells via vertical bus wires in the blue layer. Input and output wires are placed to the right and in the red layer according to the system mentioned above. This layout makes it possible to stack an optional number of flip-flops to build a counter of desired length and still have access to all inputs and outputs from the control and carry circuitry cell to the right of the flip-flops. Fig. 119 shows the complete layout.

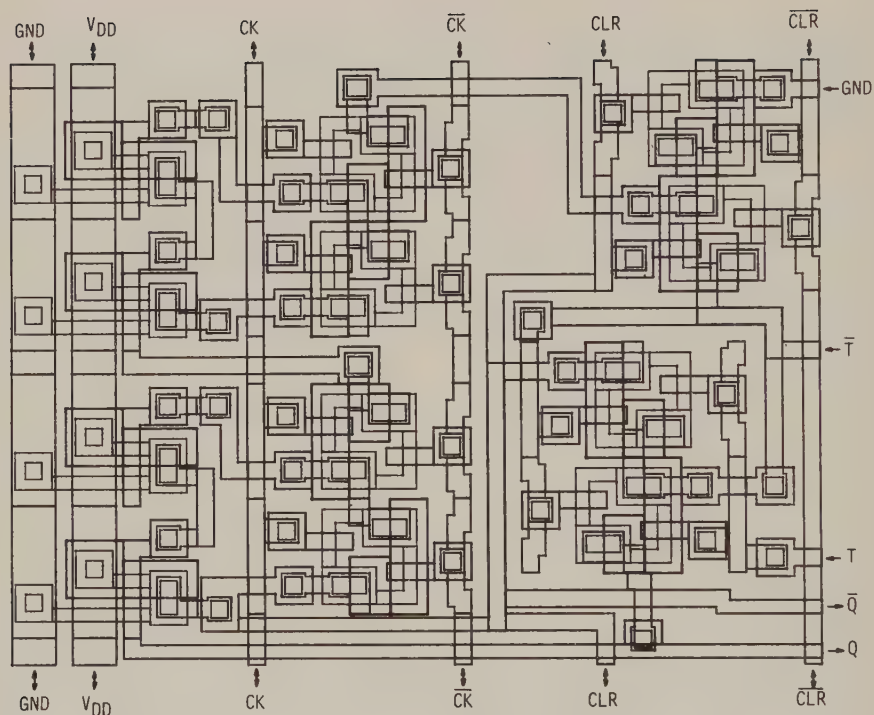


Fig. 119. Complete layout of flip-flop cell.

The remaining part to complete the synchronous decade counter in fig. 86 is the logic circuitry. The logical function can be implemented in a number of different ways. Perhaps the most systematic way is to design a switch network with the required number of inputs and outputs. In this case it would result in a cell with fairly square dimensions. Such a cell is not very favorable in this application since the connections to the four flip-flops are distributed along a comparatively long distance. Instead a solution with conventional gates is easier to adapt to a high but narrow cell. It was decided to use NAND-gates because of their faster performance. A set of cells for the required gates with 1 (inverter), 2 and 3 inputs were developed.

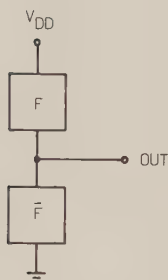


Fig. 120. Implementation of switch function in CMOS.

A general method to design a switch function F in CMOS is to use two contact networks: one pull-up network F and one pull-down network $\bar{F}$ (fig. 120). The pull-up network is always implemented with p-channel transistors since they are better conductors for high signal levels. Conversely, the pull-down network is implemented with n-channel transistors.

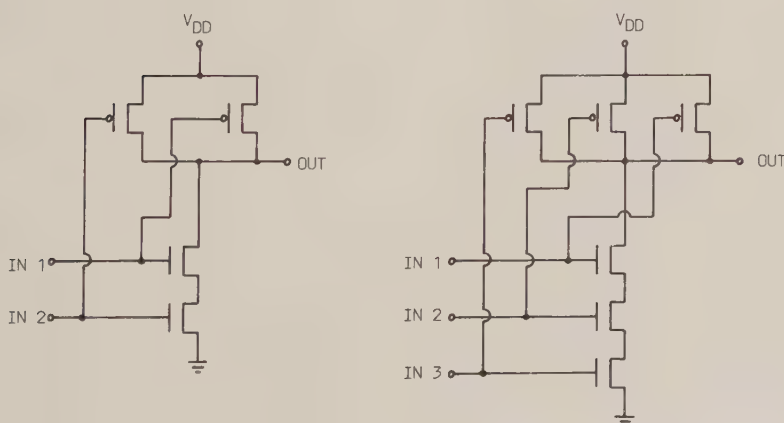


Fig. 121. Circuit diagrams for 2- and 3-input NAND-gates.

As an example consider a 2-input NAND-gate where $F = (\overline{A \cdot B})$. The Boolean expression can be rewritten as $F = \overline{A} + \overline{B}$ according to de Morgan's law. This results in a pull-up network consisting two p-channel transistors. If at least one of the inputs A or B is low, one of the transistors will be turned on and the output pulled up to a high level. The inverse function is $\overline{F} = (\overline{\overline{A} + \overline{B}}) = A \cdot B$. This implies that the pulldown network is two n-channel transistors in series. Both inputs must be high to pull the output down. Fig. 121 shows the circuit diagrams of the NAND-gates.

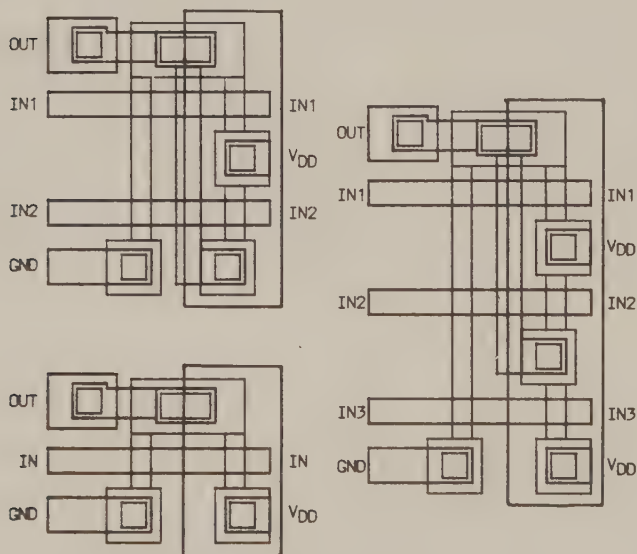


Fig. 122. Layout of the NAND-gates.

The gate cells in fig. 122 are intended to be piled on top of each other between two vertical supply leads for Vdd and ground. For that reason all cells have the same width but their heights are depending on the number of gate inputs. All inputs are accessed from both left and right side with red wires. The output is on the left side also in the red layer.

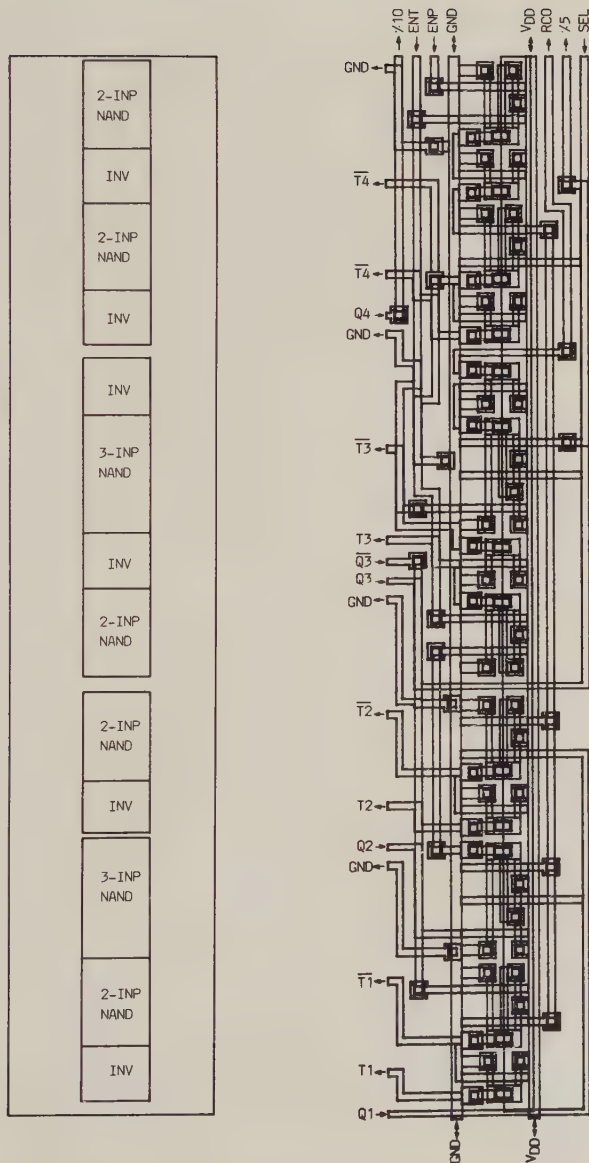


Fig. 123. Logic gate cell.

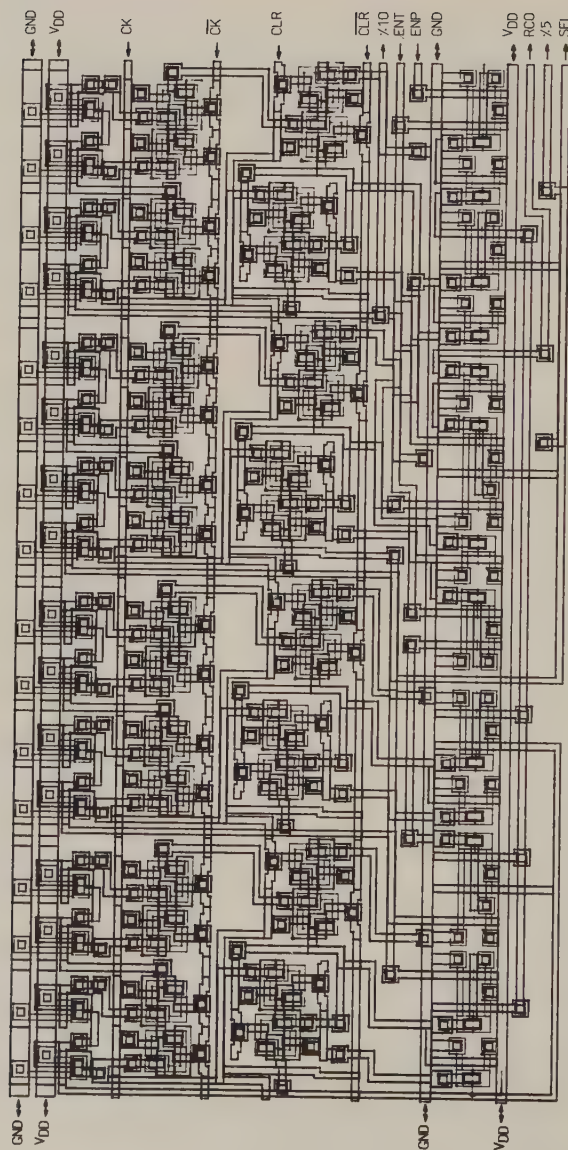


Fig. 124. Complete decade counter cell.

All p-channel transistors are situated in the right half of the cells so that a large yellow box common to all cells may be used. One corridor on each side of the gate row is used for connections between different gates and between gates and flip-flop inputs/outputs. All inputs and outputs of the decade counter are in the blue layer and placed on the top edge of the logic gate cell.

The gate network (fig. 123) exists in two versions: one which always divides by ten and one which is switchable between divide-by-five and divide-by-ten. The second type is intended for the least significant counter so that the divider ratio can be changed depending on the input frequency. The network type is selected by a parameter in the procedure call.

The next function blocks to be designed are the circuits that supply the counters with clock and reset signals. First the input 5 MHz sine wave signal has to be converted to a square wave. This is done in the input amplifier (fig. 125) which is composed of four identical minimum size inverters.



Fig. 125. Input amplifier.

Each inverter stage has a voltage gain of 10-20 which results in a total gain of at least 1000. This is sufficient to generate a square wave with adequately short rise- and fall-times even if the input amplitude is small.

An S-R latch is used to generate clock signals with the desired phase relations. Each clock signal has to drive 112 transmission gate control signal inputs. They represent a very large load capacitance which would result in a very long transition time if driven by an inverter with minimum size transistors. However, the width of the transistors can be enlarged to increase the output drive capability. Unfortunately, this also increases the input capacitance by the same factor and will slow down the input signal instead. Minimum propagation delay is obtained by using a number of buffer stages with successively greater transistor widths. The optimum number of stages n can be shown to be $\langle 44 \rangle$

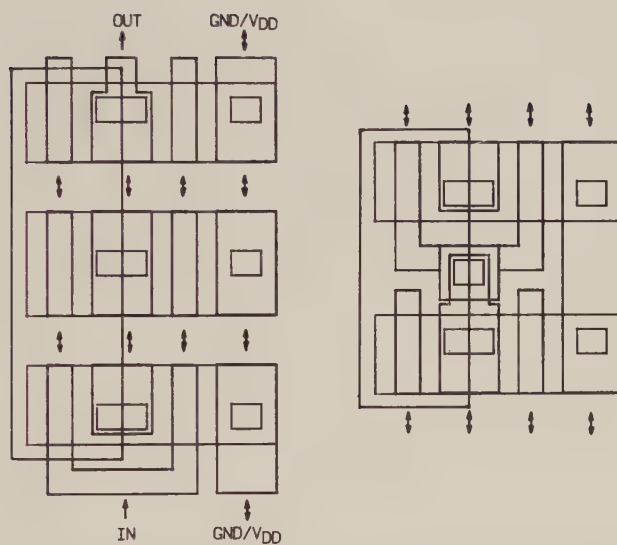


Fig. 127. Basic cells for buffer stages.

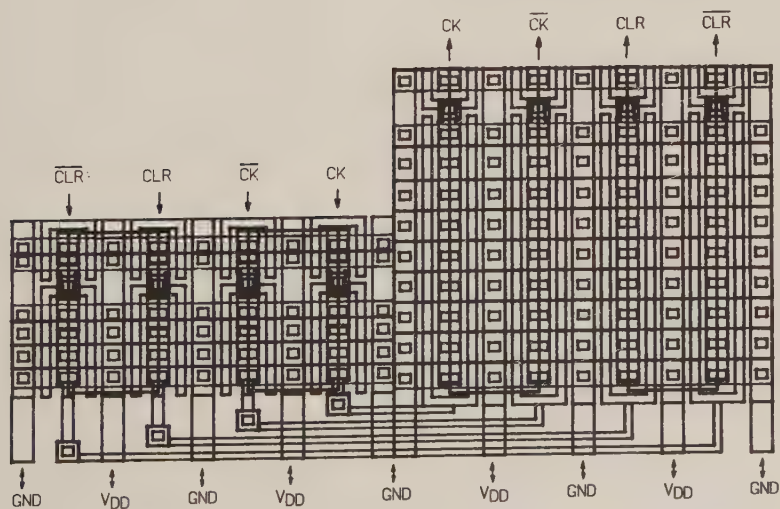


Fig. 128. Layout of the three first buffer stages.

The circuits described above (the input amplifier and the clock and clear drivers) are assembled into a larger cell with the same height as the decade counter cells. This cell (fig. 129) also contains a few other simple functions. There is a two-stage buffer for the parallel enable signal (ENP) from the least significant counter to the other six. There are also two transmission gates and an inverter to control the counting sequence of the first decade from the select input. Finally, there is a two-input NAND-gate where the carry output signal (RCO) from the most significant decade gates one period of the clock signal to the 1 PPS output.

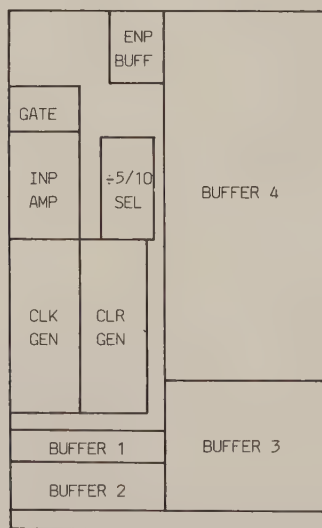


Fig. 129. Plan over the complete driver cell.

Signal inputs are placed on the left side close to the input pads. Control signals to and from the counters are placed on the top edge of the cell with access to the interconnection area.

The remaining part of the layout work is to connect the different cells to each other and to supply the necessary connections to the world outside the chip via bonding pads. A set of standard pad cells is available from a previously designed pad cell library. This library contains pads for

supply voltage connection, input pads with overvoltage protection, digital output pads with built-in buffer amplifiers, etc. These cells can be used in most standard applications and are used here for all input and output signals. It is, however, doubtful if the output pad is fast enough for the very short pulses at the 1 PPS output. The original intension was to design a special buffer to drive this output, but the shortage of time in the last phase of layout work forced us to use the standard cell.

The pad cells are located in two vertical rows on each side of the chip. All connections between the counter cells and the driver cell are brought together in the interconnection cell in the middle of the chip. The only wires outside this cell are the connections to the pad cells and the wires for supply voltage distribution. The interconnection cell contains vertical metal wires (blue) and horizontal polysilicon wires (red). The red wires are made 11 units wide to reduce the resistance and to simplify contact with to the blue wires. Normally there is no contact between these two layers since they are separated by the thick oxide. If contact is desired, a contact hole in the oxide must be made.

Many of the metal wires go straight through the cell since the upper row of counters are mirrored about the horizontal axis and placed straight above the lower row. This applies for the GND, V_{dd}, CK, $\overline{CK}$, CLR and $\overline{CLR}$. The outputs from each decade are drawn to the right or left sides of the cell and then continued outside the cell to the appropriate output pad. Fig. 130 shows the pads, the supply voltage distribution and the wiring between the driver and counter cells.

The following data applies for the complete circuit layout:

Dimensions	:	1650 x 2122	μm
Chip area	:	3.50	mm^2
Transistors	:	app. 1250	
Inputs	:	5	
Outputs	:	10	
Rectangles	:	19502	
Program lines	:	app. 1300	

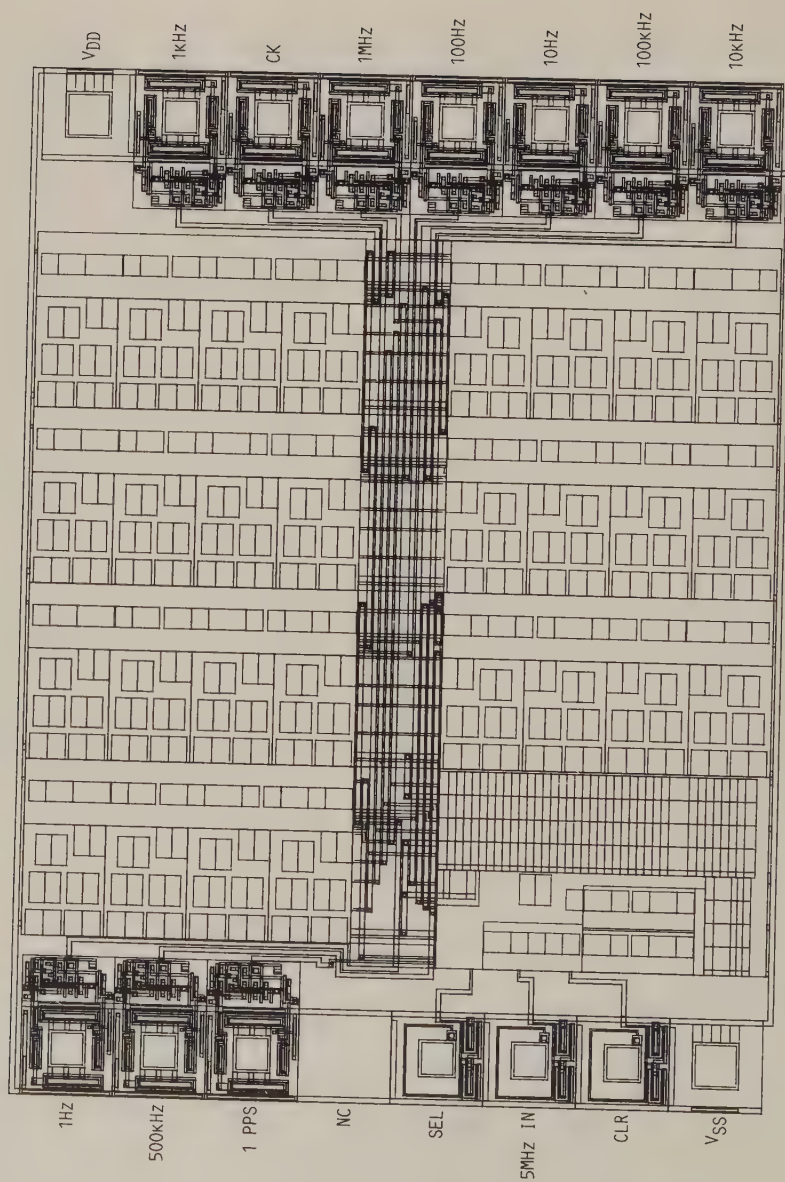


Fig. 130. Input/output pads and chip wiring.

ACKNOWLEDGEMENTS

The author wants to express his sincere gratitude to Dr. Skotte Mårtensson for the scientific guidance and for his great support and encouragement which has been a continuous inspiration in my work.

I am also indepted to Professor Lennart Stigmark for initiating this work and for the stimulating discussions regarding frequency standards.

I also want to thank the other members of the group, Mr. Bo Olsson and Mr. Göran Jönsson, for fruitful discussions and pleasant cooperation during the development of the passive hydrogen maser.

Thanks are also due to Dr. Lars Olsson for many stimulating discussions and for reading the manuscript.

For the skilful assistance with the high-vacuum system and other mechanical matters essential to the work I would like to thank Messrs. Lars Hedenstjerna and Per-Olof Ågren. This also includes Mr. Mats Ågren for producing the printed circuit boards and Mr. Bo Nilsson for carrying out the time comparison measurements.

The work has been made easier thanks to the financial and educational planning of Dr. Bertil Hansson.

I should also like to thank Miss Lena Bardskär for drawing most of the illustrations and Miss Britta Olsson for typing the parts of the manuscript where the word processor failed and for scrutinizing the text.

Dr. Ron Danielson has read the thesis thoroughly and suggested corrections of the language.

I also wish to thank Mr. Bertil Andersson, Mr. Ingvar Jonsson, Mr. Sven Mattisson, Dr. Allan Petersén, Mr. Kenneth Ström and other friends and colleagues who have contributed to the completion of this work.

REFERENCES

Frequency standards:

- <1> Blair, Byron E. :
Time and Frequency: Theory and Fundamentals
NBS Monograph 140, 1974
U.S. Government Printing Office
- <2> Kartaschoff P. :
Frequency and Time
Academic Press London, 1978
- <3> Hellwig H. :
Microwave Time and Frequency Standards
Radio Science, Vol 14, No 4, 1979, pp 561-572

Operational principles of hydrogen masers:

- <4> White H.E. :
Introduction to Atomic Spectra
McGraw-Hill 1934
- <5> Kopfermann H. :
Kernmomente
Akad. Verlagsgesellschaft, Frankfurt am Main, 1956
- <6> Lindgren I. :
Atomfysik
Universitetsförlaget, Uppsala, 1967, p 43
- <7> Mockler R.C. :
Atomic Beam Frequency Standards
Adv. in Electronic and El. Physics, Vol 15, 1961, p 8
- <8> Kleppner D., Goldenberg H.M., Ramsey N.F. :
Theory of the Hydrogen Maser
Physical Review, Vol 126, No 2, 1962, pp 603-615
- <9> Kleppner D., Berg H.C., Crampton S.B., Ramsey N.F.,
Vessot R.F.C., Peters H.E. and Vanier J. :
Hydrogen-Maser Principles and Techniques
Physical Review, Vol 138, No 4A, 1965, pp A972-983

- <10> Walls F.L. and Hellwig H. :
A New Kind of Passively Operating Hydrogen Frequency Standard
Proc. 30th Ann. Symp. on Freq. Control, 1976, pp 473-480
- <11> Walls F.L. and Howe D.A. :
A Passive Hydrogen Maser Frequency Standard
Proc. 32nd Ann. Symp. on Freq. Control, 1978, pp492-498
- <12> Howe D.A., Walls F.L., Bell H.E. and Hellwig H. :
A Small, Passively Operated Hydrogen Maser
Proc. 33rd Ann. Symp. on Freq. Control, 1979
- <13> Walls F.L. and Howe D.A. :
Timekeeping Potentials Using Passive Hydrogen Masers
Proc. 3rd Symp. on Freq. Standards and Metrology, 1981, pp 151-158
- <14> Vanier J. :
The Active Hydrogen Maser: State of the Art and Forecast
Metrologia, Vol 18, 1982, pp 173-186
- <15> Audoin C., Viennet J. and Lesage P. :
Hydrogen Maser: Active or Passive
Proc. 3rd Symp. on Freq. Standards and Metrology, 1981, pp 159-170
- <16> Zhestkova N.D. and Elkin G.A. :
Effect of a Nonuniform Magnetic Field on a Hydrogen-Maser
Translated from Izmeritelnaya Tekhnika, No 9, 1979, pp 37-38
- <17> Crampton S.B. and Wang H.T.M. :
Density-dependent Shifts of Hydrogen Maser Standards
Proc. 28th Ann. Symp. on Freq. Control, 1974, pp 355-361
- <18> Crampton S.B. :
Effects of Atomic Resonance Broadening Mechanisms on Atomic Hydrogen Maser Long Term Frequency Stability
Proc. 2nd Symp. on Freq. Standards and Metrology, 1976, pp 589-613
- <19> Desaintfuscien M., Viennet J. and Audoin C. :
Discussion of Temperature Dependence of Wall and Spin Exchange Effects in the Hydrogen Maser
Proc. 2nd Symp. on Freq. Standards and Metrology, 1976, pp 565-587

- <20> Elkin G.A. and Zhestkova N.D. :
Collision-dependent Frequency Shift in an Atomic-Hydrogen Clock
Translated from Izmeritelnaya Tekhnika, No 9, 1979, pp 36-37

Hydrogen maser design:

- <21> Menoud C., Racine J. and Kartaschoff P. :
Description et Performance Actuelles des Masers à Hydrogène du L.S.R.H.
Journal Suisse d'Horlogerie, No 7/8, 1967, pp 265-280
- <22> Hibbard L.U. :
Development of the CSIRO Hydrogen Maser
Journal of El. and Electronics Eng., Australia, Vol 1, No 3, 1981, pp 243-250
- <23> Chuang C-S. and Jair T-C. :
Atomic Time and Frequency Standards Development at Shanghai Observatory, China
IEEE Trans. on Instr. and Meas., Vol IM-29, No 3, 1980, pp 158-163
- <24> Gheorghiu O., Giurgiu L., Dragos F., Mandache C. and Bocaniciu T. :
Work on Atomic Frequency-Time Standards Based on Hydrogen Maser in Romania
Report from Central Institute of Physics, Bucharest
- <25> Gheorghiu O. and Gheorge V. :
Characteristics of a U.H.F. Discharge Used as Atomic Source for a Maser
Int. J. Electronics, Vol 39, No 3, 1975, pp 329-335
- <26> Gheorghiu O. and Mandache C. :
On a Hydrogen Dissociator for Maser
Rev. Roum. Phys., Vol 24, No 3-4, 1979, pp 317-323
- <27> Audoin C., Desaintfuscien M. and Schermann J.P. :
Phase-space Method for the Calculation of Focalization and Magnetic Separation of Hydrogen Atoms
Nuclear Instr. and Methods, No 69, 1969, pp 1-10
- <28> Becker G. and Fischer B. :
Ablenkung von Atomstrahlen in langen zylindrischen Ablenssystemen unter Berücksichtigung der Geschwindigkeitsverteilung
Zeitschrift f. ang. Physik, Vol 21, No 6, 1966, pp 492-499

- <29> Vanier J., Larouche R. :
A Comparison of the Wall Shift of TFE and FEP Teflon
Coatings in the Hydrogen Maser
Proc. 2nd Symp. on Freq. Standards and Metrology, 1976,
pp 535-564
- <30> Bell Telephone Laboratories :
Radar Systems and Components
D. van Nostrand, New York, 1949
- <31> Mattison E., Vessot R. and Levine M. :
The TE₁₁₁-mode cavity: a New, Small Hydrogen Maser
Resonator
Proc. 2nd Symp. on Freq. Standards and Metrology, 1976,
pp 615-624
- <32> Beverini N. and Vanier J. :
Tuning of Hydrogen Maser Cavity by Means of a Variable
Capacitance Diode
IEEE Trans. on Instr. and Meas., Vol IM-28, No 2, 1979,
pp 100-104
- <33> Hallén E. :
Elektricitetslära
Hallén, Stockholm, 1968
- <34> Tietze U. and Schenk C. :
Advanced Electronic Circuits
Springer-Verlag, Berlin, 1978
- <35> VAC Vacuumschmelze :
Magnetische Abschirmungen
Firmenschrift, 1969
- <36> Grinnemo L.S. :
Konstruktion av en frekvenssyntetisator avsedd för
optisk pumpning av ²³Na använd som frekvensnormal,
thesis
Lund Institute of Technology, Dept. of Applied
Electronics
IR7312, Lund, 1973
- <37> Hewlett Packard :
Step Recovery Diode Frequency Multiplier Design
Application Note 913
- <38> Gardner F.M. :
Phaselock Techniques
John Wiley and Sons, New York, 1979

- <39> Vanier J., Têtu M. and Bernier L-G. :
Transfer of Frequency Stability from an Atomic
Frequency Reference to a Quartz-Crystal Oscillator
IEEE Trans. on Instr. and Meas., Vol IM-28, No 3, 1979,
pp 188-193
- <40> Busca G. and Brandenberger H. :
Passive H-Maser
Proc. 33rd Ann. Symp. on Freq. Control, 1979, pp 563-
568
- <41> Olsson B. :
Frequency Lock and Cavity Tuning Control Systems
for a Passive Hydrogen Maser, thesis
Lund Institute of Technology, Dept. of Applied
Electronics
To be published
- <42> Mattsson T., Mårtensson S. and Olsson B. :
Fasmodulator för användning i passiv vätenormal
Lund Institute of Technology, Dept. of Applied
Electronics
LUTEDX/TETE-7003/1-16(1978), Lund, 1978

Integrated circuit design:

- <43> Mead C. and Conway L. :
Introduction to VLSI systems
Addison-Wesley, 1980
- <44> Mattison S. :
Tumregler för konstruktion i SOS/CMOS
Lund Institute of Technology, Dept. of Applied
Electronics
LUTEDX/TETE-7001/1-19(1982), Lund, 1982
- <45> Philipson L. and Breidegard B. :
Konstruktion av högintegrerade digitala kretsar
Lund Institute of Technology, Dept. of Computer
Engineering
Compendium, 1982

Measurements:

- <46> Lesage P. and Audoin C. :
Characterization and Measurement of Time and Frequency
Stability
Radio Science, Vol 14, No 4, 1979, pp 521-539

- <47> Cutler L.S. and Searle C.L. :
Some Aspects of the Theory and Measurement of Frequency
Fluctuations in Frequency Standards
Proc. IEEE, Vol 54, No 2, 1966, pp 136-154
- <48> Lesage P. and Audoin C. :
Characterization of Frequency Stability: Uncertainty
due to the Finite Number of Measurements
IEEE Trans. Instr. and Meas., Vol IM-22, No 2, 1973, pp
157-161



